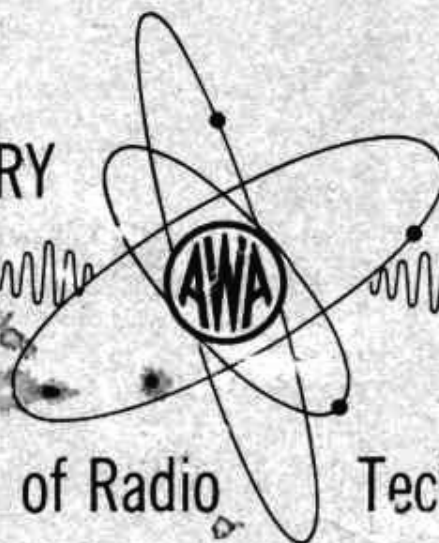


LIBRARY



of Radio Technology

MARCONI School of WIRELESS

AWA Building, 47 York Street, Sydney, N.S.W.



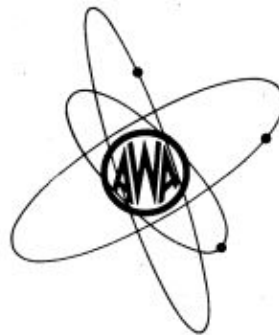
BASIC MATHEMATICS IN RADIO CIRCUITS, Stage 1 — ELECTRICAL THEORY;
Stage 2 — BASIC RADIO THEORY; Stage 3A — ADVANCED RADIO RECEIVERS;
Stage 3B — TRANSISTOR THEORY; Stage 4 — TRANSMISSION TECHNIQUES;
Stage 5 — BROADCAST STATION OPERATION; Stage 6 — RADIO AIDS TO
NAVIGATION; Stage 7 — MARINE EQUIPMENT OPERATION; TELEVISION
THEORY & PRACTICAL.

RADIO AIDS TO NAVIGATION

BOOK TWO



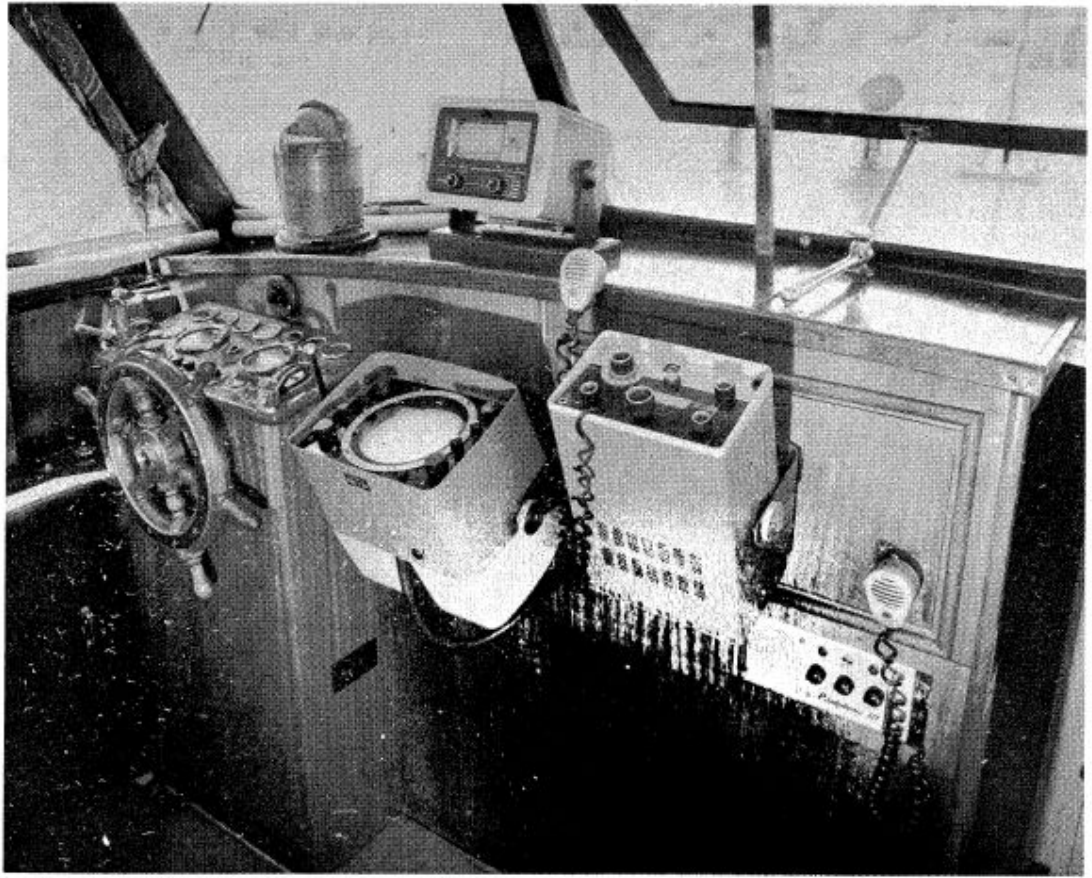
Compiled under the supervision of
The Principal, Marconi School of Wireless
in collaboration with
AWA Radio Engineering Staff.



Published by the Wireless Press for
AMALGAMATED WIRELESS (AUSTRALASIA) Ltd.
AWA Building, 47 York Street, Sydney NSW
AUSTRALIA.

(All Rights Reserved)

1967.



Comprehensive radio installation aboard motor cabin cruiser showing
(from left to right): Compass, radar, depth sounder, H. F. radiophone
and U. H. F. radiophone (frequency modulation).

· MARCONI SCHOOL OF WIRELESS
RADIO NAVIGATIONAL AIDS
CHAPTER 8. - THE MAGNETRON (Part 1.)

8-1 Introduction.

Ordinary triodes become very inefficient as generators of ultra-high frequencies, owing to transit-time effects as well as interelectrode capacitance and lead inductance. Velocity-modulated valves, such as klystrons are available for small powers only, but for very high powers the only satisfactory valve is what is known as a magnetron.

A magnetron consists of a diode operated in a strong magnetic field parallel to the length of a cathode. The electrons are emitted at right angles to the cathode, and are attracted towards the anode (plate). The magnetic field bends the electron paths around, so that the electrons may be turned back towards the cathode before reaching the anode. There is a critical value of magnetic field that cuts off the anode current for a given value of anode to cathode voltage.

In order to understand the action of the magnetron it is necessary to study the motion of electrons in a combined electric and magnetic field.

8-2 Motion in Parallel-Plane Magnetron.

Magnetrons ordinarily have the cathode and anode as concentric cylinders, but conditions are simpler in a structure in which they are arranged as parallel planes. The behaviour of cylindrical structures is similar, but differs in small details. It is easiest to start with parallel planes.

Consider a plane cathode C and an anode A arranged as in Fig. 8-1, with a positive voltage on the anode. The electric field between the two electrodes is uniform except at the edges. An electron emitted from the cathode has practically no initial velocity, but is attracted to the anode and moves in a straight line, against the lines of electric force, provided that there is no magnetic field. As the electron moves, it gains kinetic energy from the electric field.

Consider next a case when only a uniform magnetic field is applied, at right angles to the paper, say inwards, as in Fig. 8-2. Suppose that an electron happens to be moving in the space between anode and cathode, then the magnetic field applies a mechanical force to the electron, at right angles to

its motion and to the field. This force causes the electron to move around in a circle at uniform speed. Each electron path has its own centre. The size of the circle depends on the speed, but the number of revolutions per second in the circular path depends only on the magnetic field, being given by the formula $f = 2.8 \times 10^6 B$, where B is the intensity of the magnetic field in lines per square centimetre. The value of f is known as the "cyclotron frequency"; one type of magnetron oscillates at its cyclotron frequency.

In combined electric and magnetic fields at right angles, as in Fig. 8-3, the electrons have a motion that can be regarded as a steady drift parallel to the surface of either electrode, at right angles to both fields, with a superimposed motion in a circle. The drift velocity is given by the formula
$$v = \frac{10^8 E}{B}$$

where E is the electric field strength in volts per centimetre, and v is in cm. per second. The circular motion is at f revolutions per second, where f is the same as in the absence of an electric field, and the circle is of such a size that the speed of motion around it is equal to the drift velocity.

The total motion is like that of a point on the rim of a wheel rolling along the cathode. The point at the bottom of the wheel has a forward component of velocity due to the drift, and a backward component due to the rolling, which is equal and opposite to the forward component. The electron therefore starts from rest at the cathode. As the wheel rolls along there is a difference in direction of the component velocities, and the resultant has both vertical and horizontal components. At the summit, corresponding to a 180° rotation of the wheel, both components are in the same direction, and the instantaneous velocity is twice the drift velocity. The electron moves in a curve known as a cycloid (named from the rolling wheel), consisting of a series of arches with sharp points known as cusps. The horizontal component of velocity varies from twice the drift velocity to zero, and the average horizontal velocity is equal to the drift velocity. At the cusps the electron is instantaneously at rest, with neither vertical nor horizontal velocity.

Each arch of the cycloid corresponds to one revolution of the rolling wheel, during which the direction of resultant motion turns through 90° , from vertically upwards to horizontal, in the first half revolution, and through a further 90° , from horizontal to vertically downwards, in the second half. The direction of motion changes uniformly with the time, at exactly half the rate at which the wheel turns. The speed can be proved to be proportional to the square root of the distance from the cathode, and is zero at the cusps, which are at the cathode surface.

If the rolling circle has a diameter greater than the distance between anode and cathode, the path of the electron will cut the anode surface. This means that the electron will strike the anode and be collected. Otherwise it will return to the cathode, which it just reaches as it comes to rest, and can hardly be said to strike.

In a cylindrical structure, the electron paths can be obtained by rolling a circle around the cathode. Fig. 8-4 shows several paths, for different values

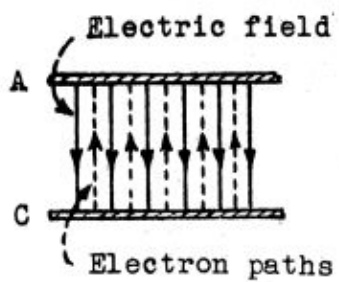


Fig. 8-1

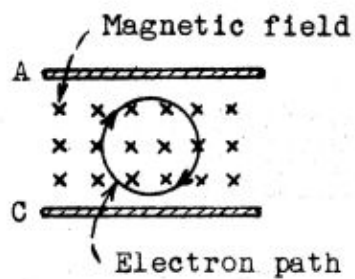


Fig. 8-2

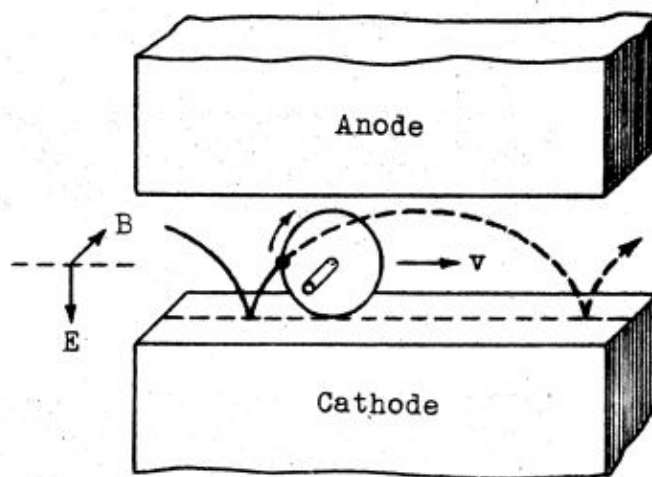


Fig. 8-3

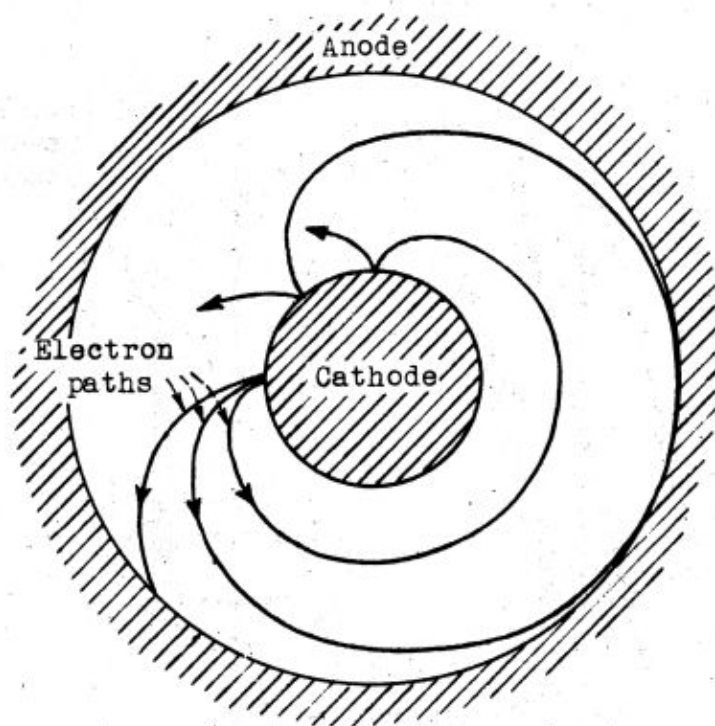


Fig. 8-4

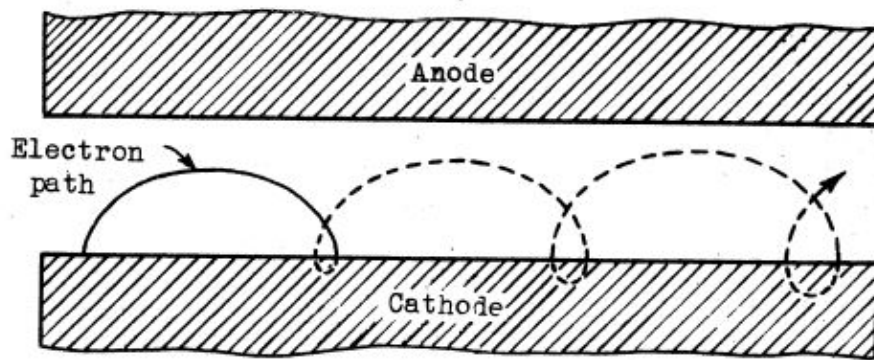


Fig. 8-5

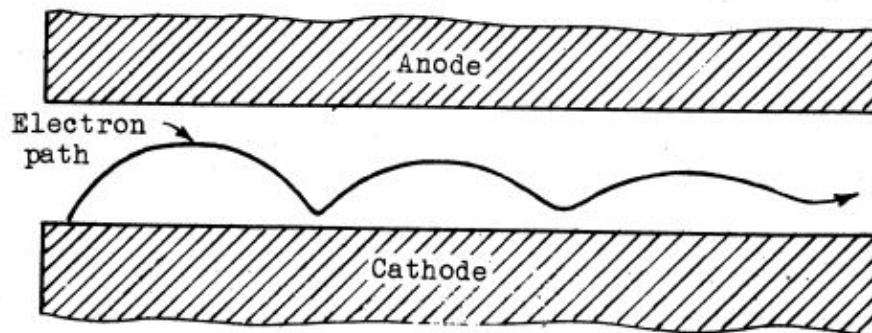
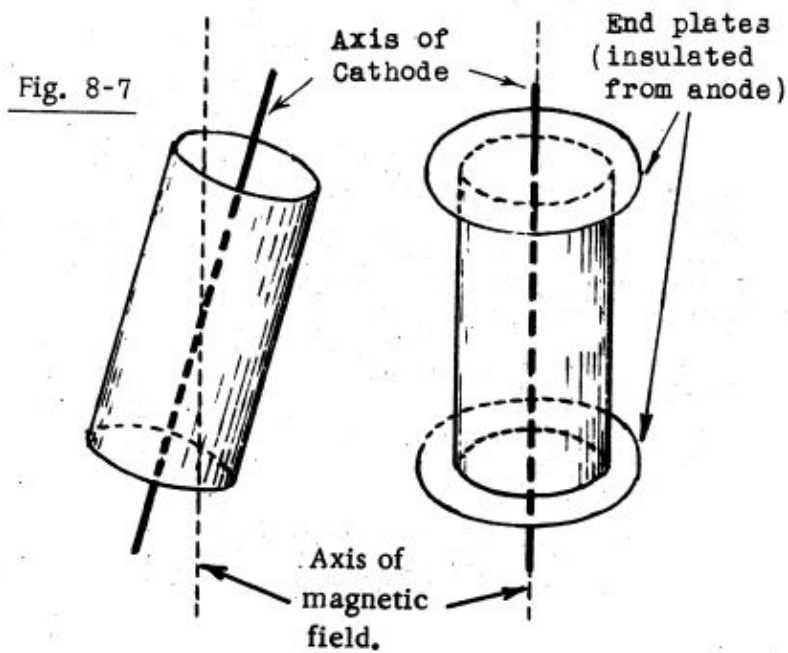


Fig. 8-6



(a) Magnetron with axis tilted with respect to magnetic field.

(b) Magnetron with end plates.

of magnetic field.

From the above theory, all the electrons emitted should reach the anode if the magnetic field has less than a certain critical intensity, and none if it is greater. In practice the cut-off is gradual, because the electrons are emitted with small velocities which vary over a certain range, and because of the effects of space charge, which sets up a complex electric field.

8-3 TYPES OF MAGNETRON OSCILLATORS.

There are three different types of magnetron oscillators, of which the first two are of more historical interest, as the third type is so much more efficient that it has superseded them. Still, the earlier types are worth describing for the manner in which they illustrate principles that still apply in the third type.

8-4 Cyclotron Frequency Oscillator.

If a parallel-tuned circuit, resonant at the cyclotron frequency, is connected in series with the HT supply between the anode and cathode, oscillations can take place. An r-f current flows through the magnetron and builds up an r-f voltage across the tuned circuit, so that the total voltage between anode and cathode varies alternately above and below the applied d-c voltage. The half-cycle during which the voltage is above the d-c value will be called positive. The magnetic field is usually chosen just beyond the cut-off value, and no anode current flows unless oscillation starts.

For simplicity we again take a parallel-plane structure. The voltage between anode and cathode sets up a uniform electric field, the lines of force being directed vertically downwards from anode to cathode. This field may be resolved into a d-c component and an r-f component. On the positive half-cycle these assist each other, and are both downwards, but in the negative half-cycle the r-f field is reversed. It is then directed upwards, and opposes the d-c field.

Consider an electron that starts out from the cathode at the beginning of the positive half-cycle. It is attracted by the anode, and gains energy from both the d-c and r-f fields, which are acting together. The path of the electron is shown by the full-line curve in Fig. 8-5. The magnetic field, at right angles to the paper, bends the path around, so that at a certain distance from the cathode the electron is moving horizontally. This distance depends on the anode voltage, and is therefore greater during the positive half-cycle. With the d-c component of voltage alone the electron would pass close to the anode, but if the r-f voltage is large enough the electron may strike the anode and be collected. The energy that it has acquired from both d-c and r-f sources is then dissipated on the anode.

If the electron fails to strike the anode it will move back towards the cathode. It requires one half-cycle to turn the direction of motion through a right angle, so that the electron is returning towards the cathode in the negative half-cycle. It is assisted in its return by the r-f field, which has reversed, but hindered by the d-c field, which is stronger. The electron therefore continues to

draw energy from the r-f field, but now restores energy to the d-c field, and reduces its own store of energy. When the electron reaches the cathode surface again, it has returned all the energy that it drew from the d-c field in the outward motion, but still has the energy drawn from the r-f field. This energy is given up to the cathode when the electron strikes it. The net result is that an electron emitted at the time stated opposes the oscillation by drawing energy from the r-f field and dissipating it at the anode or cathode. Such an electron is said to be unfavourable to the oscillation, and may be referred to as an unfavourable electron.

If the electron were able to continue its motion, instead of being stopped by the cathode, the path would be as shown dotted. The magnetic field would continue to bend the path around, and a small loop would be formed. The size of this loop depends on the amount of energy that the electron has acquired from the r-f field. As it takes some time to traverse the part of the loop below the cathode surface, the electron would emerge as a re-emitted electron later in the next positive half-cycle.

An electron emitted in any part of the positive half-cycle is still unfavourable, because on the average more energy will be drawn from the r-f field than will be restored to it. The re-emitted electron, unless it struck the anode, would again be brought back to the cathode and travel around a larger loop. The effect of unfavourable electrons is then to build up the circular part of the motion at the expense of the r-f field.

Actually all such electrons are collected during the first cycle and removed from the space between cathode and anode, so that they can act unfavourably during this cycle only. Consider now an electron emitted at the beginning of the negative half-cycle. In this case the electron gains energy from the d-c field, which attracts it, but loses energy to the r-f field, which tends to oppose its motion. After a half-cycle the electron is moving horizontally, and certainly misses the anode. In the next half-cycle it returns towards the cathode, as the path has been bent around by the magnetic field. It is now restoring energy to the d-c field, but continues to deliver energy to the r-f field, which has now reversed and opposes its return to the cathode. The result is that the vertical component of velocity is lost before the electron reaches the cathode, as shown in Fig. 8-6. At the end of the cycle the electron is again moving horizontally, but farther from the cathode. The variations in distance from the cathode during a cycle become less and less in each cycle, until after about four or five cycles the electron is moving nearly horizontally all the time.

The electrons may be removed by either of the arrangements shown in Fig. 8-7. If the magnetic field is arranged at an angle to the axis of the cathode, the electrons will have a velocity component parallel to the magnetic field and unaffected by it. This component itself has a component parallel to the axis of the cathode. The electrons therefore drift along beyond the end of the anode. Alternatively, circular end-plates can be mounted so that an electric field can be applied between them, again producing an electron velocity component parallel to the axis of the cathode. The electrons are then collected on the more

RADIO NAVIGATIONAL AIDS.

positive end-plate. Both of the end-plates must be insulated from the anode. It is better to use end-plates even with a tilted magnetic field, as otherwise the electrons charge up the glass walls. It is convenient to connect one end-plate to the cathode, the other to the HT supply, without including the tank circuit.

Equal numbers of electrons are emitted at favourable and unfavourable times. But the favourable electrons stay in the field for several cycles, while the unfavourable electrons are removed during the first cycle. Thus, on the average, more energy is delivered to the r-f field than is withdrawn from it; and the magnetron is able to maintain an oscillation.

The power relations in the magnetron are rather complex. The power input is equal to the HT voltage multiplied by the d-c component of net cathode current, which is the emission current minus the return current of unfavourable electrons collected at the cathode.

Of the emitted electrons, half are favourable and eventually reach the positive collector plate at the drift velocity. Most of their energy has been converted to r-f, and the dissipation at the collector is not large. Some of the unfavourable electrons reach the anode and dissipate there both d-c and r-f power. The others return to the cathode, where they have lost their d-c power and dissipate only r-f power. The balance of the r-f power is available for the tank circuit.

The r-f component of anode current is produced by induced charges as the favourable electrons approach and recede from the anode.

8-5 Negative-resistance Magnetron Oscillator.

Another type of magnetron has the cylindrical anode split lengthwise into two sections, and a lead is brought out from each, so that the anode potentials may be separate. The magnetic field is again applied parallel to the axis.

Consider the static characteristics of the valve. If the two halves of the anode are at the same potential, the splitting has no effect. If the magnetic field is strong enough, there is no current to either anode section. But when the anodes are at different potentials, an additional electric field between them is superimposed on the original radial electric field, and the resultant field is complex, having lines of force as shown by the thin lines of Fig. 8-8. In such a field the electrons are not returned to the cathode when their paths are bent back by the magnetic field, but move through a number of loops and finally strike the anode of lower potential. Figs. 8-8 (a) and (b) show paths for electrons leaving the cathode on opposite sides, and it is seen that both paths end on the lower anode section, which is relatively negative to the other.

If the current I_1 to one of the anodes is plotted against the potential difference $E_1 - E_2$ between the two anodes, keeping the average of the two potentials constant, the curve obtained will slope downwards, and so correspond to a "negative" resistance, such as occurs in screen-grid and dynatron valves. This is the reason for using the term negative-resistance magnetron.

If a tuned circuit is connected between the two anode sections, as in Fig. 8-9, oscillations will take place at the natural frequency of the tuned circuit, in the same way as in a dynatron. The operating frequency must be chosen well

below the cyclotron frequency, as the period must be long enough to allow the orbits of Fig. 8-8 to be completed.

8-6 TRAVELLING-WAVE MAGNETRON.

The modern type of magnetron, which is known as the travelling-wave magnetron, is capable of very high output, at good efficiency, and has therefore superseded the earlier types. It has a central cathode which is surrounded by an anode divided into several segments, always an even number. These are connected together by cavity resonators instead of external tuned circuits. Physically it is a development of the split-anode magnetron, but the operation is quite different.

Magnetrons by different manufacturers vary greatly in external appearance, but they have certain features in common. The active parts are completely enclosed by metal, and cannot be seen. There are two heater leads, to one of which the cathode is connected. They are usually protected by pyrex glass shrouds, which are sometimes further protected by a wide metal tube. The output lead is either opposite the heater leads or at right angles. The output usually begins in coaxial form, which may be retained, or alternatively a change may be made to a wave guide, so that the output appears as a short wave-guide section with a mounting flange.

The anode block is surrounded by cooling fins, to dissipate the heat produced. In large sizes forced air cooling is used, a blast of air being produced by a fan.

In some magnetrons, the magnet is made an integral part of the assembly, while in others the magnetron is slid in between the poles of the magnet and clamped to the yoke.

Fig. 8-10 shows a cut-away view of one type of magnetron with the top cover removed. The purpose of the connecting wires between alternate segments is explained later.

Fig. 8-11 shows the anode block of a magnetron, with portion cut away to show the cathode. This is a thick cylinder with end-plates, and is oxide-coated to provide an emitting surface. A heater is provided inside the cathode. The anode is machined out of a solid copper block. Each resonator is formed of a cylindrical hole, with a parallel-walled slot extending to the central cylindrical surface. The anode segments are thus formed between the slots. The space between the cathode and the ring of segments is called the interaction space. This serves to couple the resonators together and to the electron motions. Power is taken out by coupling a loop to one resonator, and leading away by a coaxial line.

The operating frequency is very close to the resonant frequency of each resonator, being slightly modified by the effects of coupling between the resonators. Each resonator then has a high impedance at the mouth. As regards dc, the potential of the whole anode block is uniform, but r-f differences of potential can develop between the anode segments, and so set up an r-f field in the interaction space as well as in the resonators.

RADIO NAVIGATIONAL AIDS.

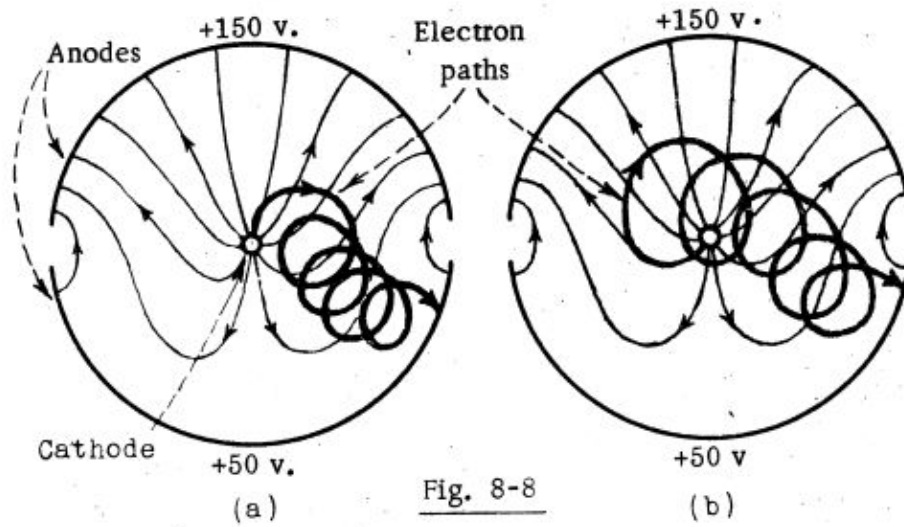


Fig. 8-8

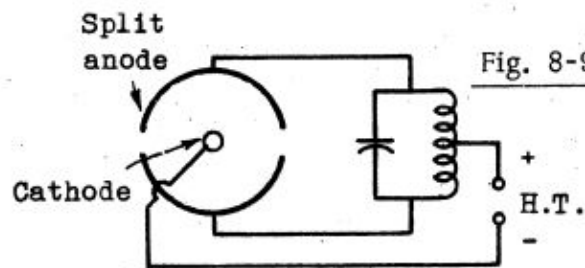


Fig. 8-9

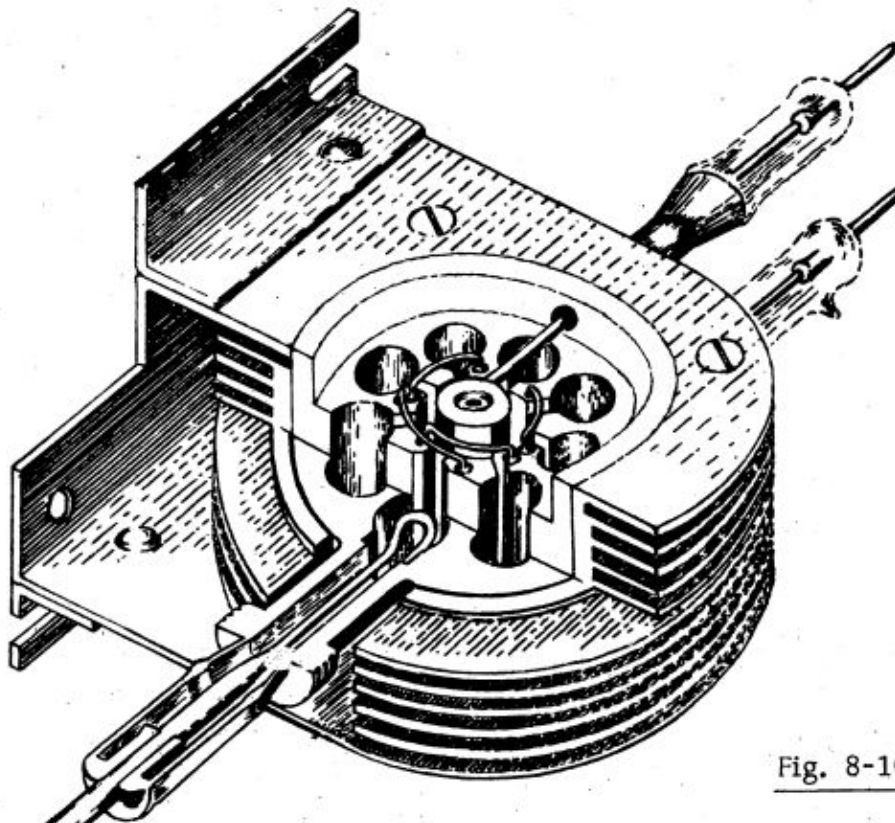


Fig. 8-10

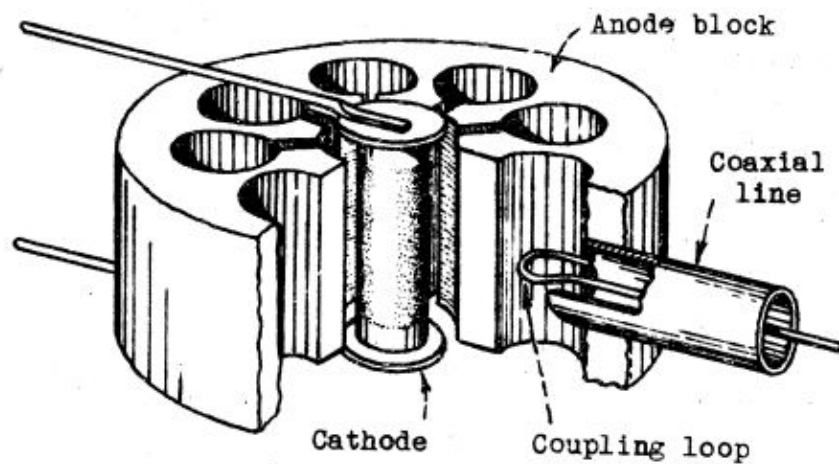


Fig. 8-11

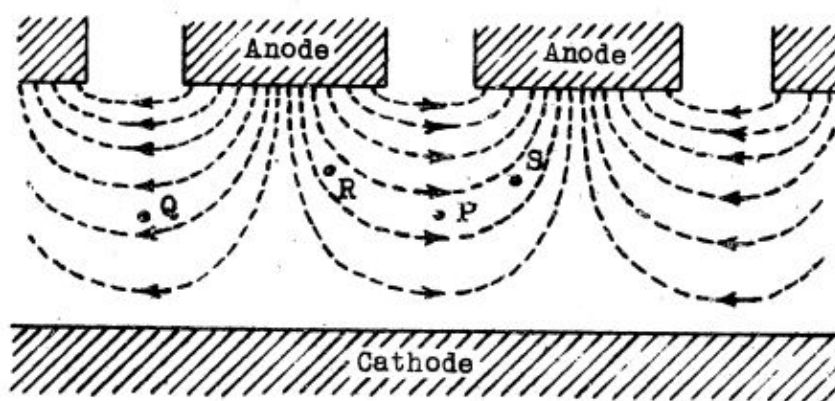
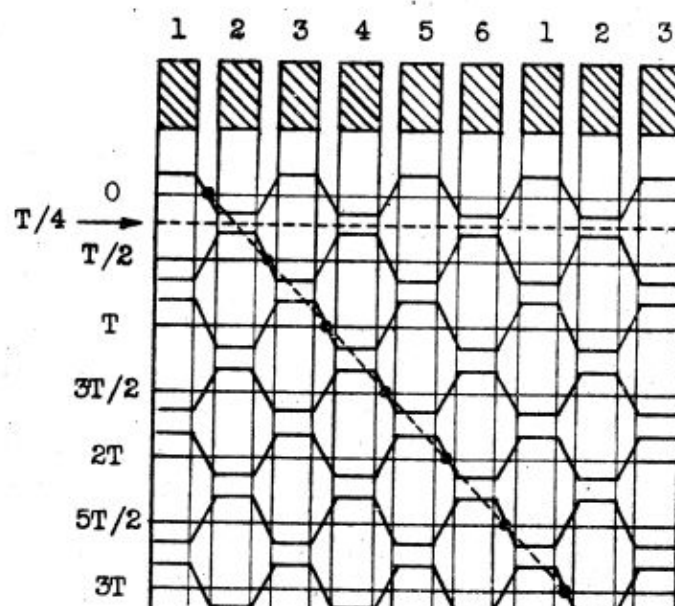


Fig. 8-12



RADIO NAVIGATIONAL AIDS.

The anode block is earthed in operation, and negative H. T. is applied in pulses to the cathode. The peak power obtainable is very large in comparison with the allowable steady dissipation, and so the magnetron is well adapted to pulse operation.

8-7 Formation of R-F Field.

The electrons emitted by the cathode are attracted towards the anode, but their paths are bent by the action of the magnetic field, so that they have a velocity component around the cathode and can drift past several anode segments and gaps. By means of a mechanism described later, the electrons are gathered together in a number of groups, leaving certain regions practically free of space charge. Each group induces a positive charge on the anode segments as it passes. This positive charge flows across the width of a segment, around the boundary of a resonator, across the width of the next segment, and so on, keeping pace with the inducing space charge cloud. The adjoining segments carry negative charges, the algebraic sum of the charges being zero. At any point on a segment, the charge density becomes alternately positive and negative, varying at the operating frequency. The potential is uniform across the width of each segment, but there are strong electric fields between segments, across the mouths of the resonators. At certain points in the cycle the electric field across the mouth disappears, when the charge densities momentarily pass through zero, one becoming negative and the next positive. At such instants the charge is flowing most rapidly around the back of the resonators, in which strong magnetic fields are produced. Although this is not the only possibility, the magnetron is usually so operated that alternate segments are in exact phase opposition, so that when the first, third, etc. segments are positive the second, fourth, etc. are negative to the same extent. All the odd-numbered segments are at the same potential, and all the even-numbered segments are at another potential. The fields across successive gaps are then equal but opposite in phase and consequently in direction. The field extends some distance into the interaction space as in Fig. 8-12, which is drawn for a parallel-plane structure for the sake of simplicity.

The r-f field has both horizontal and vertical components, both of which are important. The horizontal field is concerned with the maintenance of the oscillation, while the vertical field is responsible for keeping the electrons in groups. Both components contribute towards removing unfavourable electrons.

8-8 Maintenance of Oscillation.

The electrons emitted are speeded up or accelerated by the radial d-c electric field, and develop a horizontal drift velocity owing to the magnetic field. This is assumed to be towards the right.

An electron moving horizontally is acted upon by the horizontal component of the r-f field. The lines of force run from positive to negative segments, and electrons are attracted towards the positive segments. Thus the force applied to an electron is opposite in direction to the lines of force. An electron at the point P in Fig. 8-12 is being retarded or slowed down by the r-f field, and is

losing energy, which is delivered to the r-f field and goes to increase the energy stored in the oscillating system, which includes the resonators and the interaction space.

Suppose that the relative values of the HT voltage and the magnetic field are so adjusted that an electron drifts from the neighbourhood of one gap to the next gap in half a cycle, passing the centres of the gaps at the instants when the horizontal electric field has its peak retarding value, and the potential curve is falling most steeply. The horizontal r-f field is small while the electron crosses in front of the segment, because this occurs at about a quarter of a cycle from the peak, and because the field is in any case mainly vertical. During this time the electron gains energy from the d-c field.

When the electron crosses the next gap the field has reversed, and so the electron again loses energy to the r-f field. It can pass several gaps and make several contributions to the r-f field. Such an electron is in a favourable position to maintain the oscillation.

The periodical slowing down of the electron reduces the effect of the magnetic field upon it, so that it eventually reaches the anode and is collected. Owing to the cylindrical structure, the electron would have to drift faster at greater radii to pass from gap to gap once per half-cycle, so that there is a tendency to get out of step and become unfavourable. Before the electron can actually become unfavourable it is collected and removed. There is no need of any special provision for removing the electrons.

Consider now an electron at the point Q; this is in an accelerating field, and is speeded up. The magnetic field then has a greater effect and tends to return the electron to the cathode. The extra energy of the electron is obtained from the r-f field, so that the electron is unfavourable and needs to be removed. It is always removed at the cathode during the first loop of its path.

Equal numbers of favourable and unfavourable electrons are emitted, but the favourable electrons stay longer in the interaction space, and consequently most of the electrons present at any time are favourable. Thus, on the average, energy is being delivered to the r-f field, which builds up until the r-f power dissipated in the resonators or output line or transferred to the load is equal to the power supplied to the r-f field.

8-9 Anode Potential Wave.

The repeated interaction of the r-f field and the electrons can be made clearer by means of Fig. 8-13, which represents the variations of potential of the anode segments at different times. This is a stationary wave pattern with a node at the centre of each slot. The horizontal field however depends on the slope. The period of the r-f oscillations is denoted by T . At the top are shown six anode segments and the first three repeated, so that every region is shown at least once without a break. Below this is shown the anode potential wave at the instant when the first segment is maximum positive. The dot represents an electron at the middle of the first gap, and the downward slope of the wave indicates a retarding field. A quarter of a cycle later the r-f electric field is zero every-

RADIO NAVIGATIONAL AIDS.

where, as shown dotted, and the electron is then passing the second segment. The field now reverses, and at half a cycle from the beginning the second segment is maximum positive. If the electron is now in the centre of the second gap, it will again be in a retarding field. After another half cycle the electron is in third gap, and after a total of three cycles it has progressed completely around the anode to the first gap again. The r-f field is zero while the electron crosses any segment, and retarding while it crosses any gap. The inclined dotted line joining the dots indicates the average drift velocity, which must be close to the correct value if the electron is to remain favourable.

8-10 Electron Paths,

At any point, the r-f field is alternating. But if we consider the horizontal field acting on an electron drifting from gap to gap in a half-cycle, this will have a definite average value, with superimposed variations. If these are neglected, we have an approximate condition of a uniform field with vertical and horizontal components, giving an inclined resultant, as shown in Fig. 8-14, which is drawn for the case of an accelerating r-f field, and so applies to unfavourable electrons. The resultant field slopes downwards and to the left, and a similar field could be produced by a pair of imaginary parallel electrodes with their faces at right angles to the resultant, as shown by the dotted lines. The resultant field being uniform, the electron paths will be cycloids with base parallel to the imaginary electrodes, as in Fig. 8-15. The path cuts the cathode surface in the first arch and the unfavourable electron is removed, but the extension of the path is shown dotted to indicate the real nature of the path.

When the r-f field is retarding, the resultant field is downwards and to the right, and the base of the cycloidal paths slopes upwards, as in Fig. 8-16. The electrons can then pass through several arches and finally reach the anode. At each cusp the electron has lost all its energy, which has been delivered to the r-f field, and when it reaches the anode the electron has only the energy gained during the last incomplete arch. If we keep the ratio of HT to magnetic field constant the drift speed is fixed, and then the magnetic field may be set at such a value that a cusp occurs near the anode. The loss of energy at the anode will then be small.

8-11 Formation of Electron Groups.

The r-f field may be approximately described as vertical under the anode segments and horizontal under the gaps. The retarding field in the gaps acts as described above, and opposite the segments the vertical field acts to form bunches or groups of electrons. When the r-f vertical field assists the d-c field the drift velocity is increased, when it opposes the d-c field the drift velocity is decreased.

Consider an electron located at the point P in Fig. 8-12 at the moment when the retarding field at the gap is maximum. A little earlier it was under the action of the downward vertical component of the r-f field, so that the total vertical field was increased. This increased the drift velocity above the average value. A little later, when the electron is under the next segment, the vertical component of r-f field is reversed, and the drift speed is reduced. On the

average the vertical component of the r-f field has practically no effect.

Consider now an electron located at the point R when the r-f field is maximum. This electron will be speeded up in the same way as the other. But when it passes under the next segment the wave is well past the peak, and the electron will not be slowed down so much. It will therefore tend to overtake the first electron.

In the same way an electron located at the point S when the r-f field is maximum will be speeded up slightly when under the previous segment, but slowed down considerably under the segment where it is shown. The result is that it falls back to be with the others. The electrons tend to collect in groups that at the moment when the field has its peak value pass the centres of gaps where the field is retarding.

In the field shown in Fig. 8-12, a group will form at the middle gap, and at every alternate gap around the anode. The bunching is entirely separate from the interaction with the retarding field.

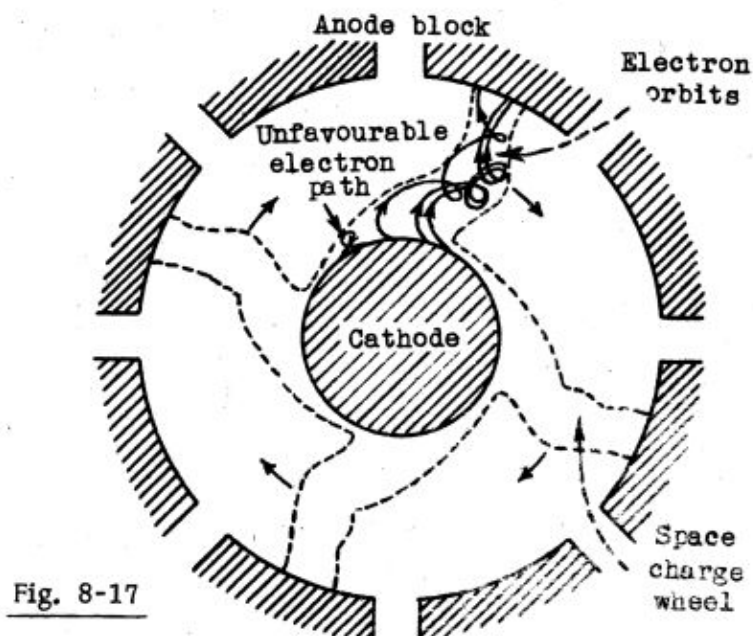
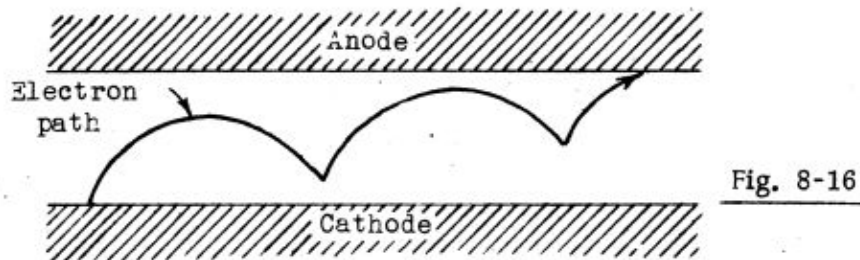
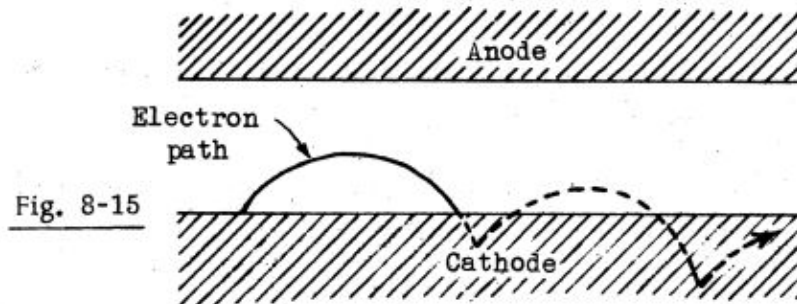
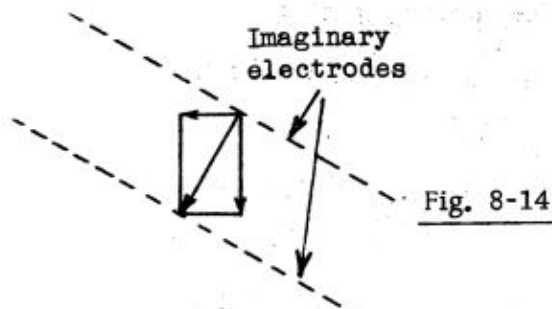
8-12 Space Charge Wheel.

The electrons are cleared from in front of segments with accelerating fields, and collect in groups in front of those with retarding fields. The space charge collects in the shape of wheel with a number of spokes, as shown in Fig. 8-17 for the case of eight segments. Four spokes are shown, corresponding to half the number of segments. Each spoke rotates from one segment to the next in half a cycle, and excites a resonator as it sweeps past. This gives a very simple and clear picture of the maintenance of the oscillation by the space charge.

The hub of the wheel, which extends to only a short distance from the cathode, consists of unfavourable electrons which are turned back to the cathode before reaching out far into the interaction space, one such path being shown. The favourable electrons are in the spokes or in the part of the hub from which these spring.

The electrons in the spokes move out from cathode to anode along the complicated orbits shown. These are relative to the spokes. For the motion relative to a fixed point the motion of the spokes must be added in, and then the loops shrink to cusps. Each arch of the cycloidal path then obtained corresponds to one cycle at the cyclotron frequency, which has no definite relation to the operating frequency.

RADIO NAVIGATIONAL AIDS.



QUESTION PAPER FOR CHAPTER 8.

1. Explain the conditions which make an electron favourable or unfavourable in the travelling wave magnetron.
2. Describe the mechanism for the production of bunches of electrons in a travelling-wave magnetron.
3. Describe the motion of an electron in electric and magnetic fields, separately and combined at right angles.
4. Describe the operation of the travelling-wave magnetron in terms of the space charge wheel.
5. Compare the methods of removal of used favourable electrons in the older and newer types of magnetron.

MARCONI SCHOOL OF WIRELESS
RADIO NAVIGATIONAL AIDS.
CHAPTER 9. THE MAGNETRON (Part 2).

9-1 MODES OF OSCILLATION.

Although successive resonators are usually operated in phase opposition, there are other possibilities, corresponding to different modes of oscillation. These must also be considered in order to understand the special constructional features and operating conditions required to ensure that oscillation will occur in the desired mode only.

It is evident that the phase differences between the oscillations in successive resonators, and the phase differences between the potentials of successive anode segments, must all be equal. It is also evident that the sum of the phase differences between successive segments must be a whole number of cycles, and similarly for the resonators.

Thus in a six-segment anode the phase difference between successive segments must be a multiple of 60° . Since any whole number of cycles may be discarded, this gives the six possible values 60° , 120° , 180° , 240° , 300° , and 360° or 0° , of which the 180° mode is that previously described. Since 300° is equivalent to -60° , the first and fifth modes differ only in sense around the anode, with the space charge wheel rotating in opposite directions, and similarly with the 120° and 240° modes. In each case we may consider that there are two components of a "double" mode. Thus there are four different modes, and these correspond to four frequencies of oscillation.

If we regard the resonator slots as condensers and the holes as inductances, the equivalent circuit of the resonator system shown in Fig. 9-1 corresponds to six mutually coupled tuned circuits. Such an arrangement can resonate in several modes at different frequencies. These frequencies depend not only on the resonant frequency of each resonator alone, but also on the closeness of coupling and the phase difference. This may be seen in the simple case of two identical coupled circuits, where the effective inductance takes the two values $L + M$, and $L - M$, corresponding to phase differences of 0° and 180° respectively between the currents in the two circuits.

The successive phase differences indicate progressive or travelling waves, having speeds such that they go around the anode once in a whole number of

cycles. The wave forms are of complex shape, being composed of horizontal and inclined straight lines, but are periodic waves in space. As such they must be composed of a number of "space" harmonics, although all components have the same frequency. Consider the case of 60° phase difference, so that the main or fundamental component of the wave travels around the anode once per cycle. With this wave we can associate other components in which the phase difference is 60° plus any whole number of cycles, so that these components travel around the anode in 7, 13, or 19 cycles, and so on, corresponding to $1\frac{1}{6}$, $2\frac{1}{6}$, $3\frac{1}{6}$ cycles per segment. Such harmonics build up the actual waveform. The different components travel at different speeds, having different wavelengths, but the same frequency, and the resultant wave changes in shape as it moves. But after every sixth of a cycle of the fundamental the original shape reappears, having moved on one segment.

Of the above components, only the fundamental extends deeply into the interaction space, the others being confined to the immediate neighbourhood of the anode surface. Thus in the middle of the interaction space the waveform is a pure sine wave.

Fig. 9-2 shows the anode waveform at intervals of a twelfth of a cycle, over one complete cycle. The two waveforms shown are the extreme limits of variation. At the times corresponding to even sixths of a cycle, the retarding field has its peak value in turn in each of the slots, moving from slot to slot each sixth of a cycle. The dots represent an electron located in these same slots in turn at the corresponding instants, and therefore drifting to the right at the same speed as the wave. Whenever it is passing a slot it is in a retarding field, but when it passes a segment the horizontal field is zero. Thus the electron delivers energy to the travelling wave. In this mode an electron passes six slots per cycle, compared with two in the mode with 180° phase difference. The drift speed is therefore greater, and a higher value of HT is required to excite the mode if the magnetic field is held constant.

The same mode can be excited by a stationary wave, if the electrons move across three segments per half cycle as indicated in Fig. 9-3. In this case there are potential nodes at two opposite gaps, and the resonators are unequally excited. The output would be taken from a resonator at a potential node, because the r-f field across the gap is greatest.

9-2 Mode Fields in Cylindrical Structure.

Some idea of the r-f fields in a six-segment magnetron in the different mode is given in Figs. 9-4, 9-5, and 9-6, which refer respectively to the cases of successive phase differences of 60° , 120° , and 180° . In each case the segments are labelled with their phase angles for the moment when the top segment is at phase zero, corresponding to the positive peak. The potential distribution at the centres of the segments then follows a cosine wave. The electron drift is supposed to be clockwise.

In the first mode, shown in Fig. 9-4 the greatest voltage difference is between the top and bottom segments. The field is mainly vertical and is

RADIO NAVIGATIONAL AIDS

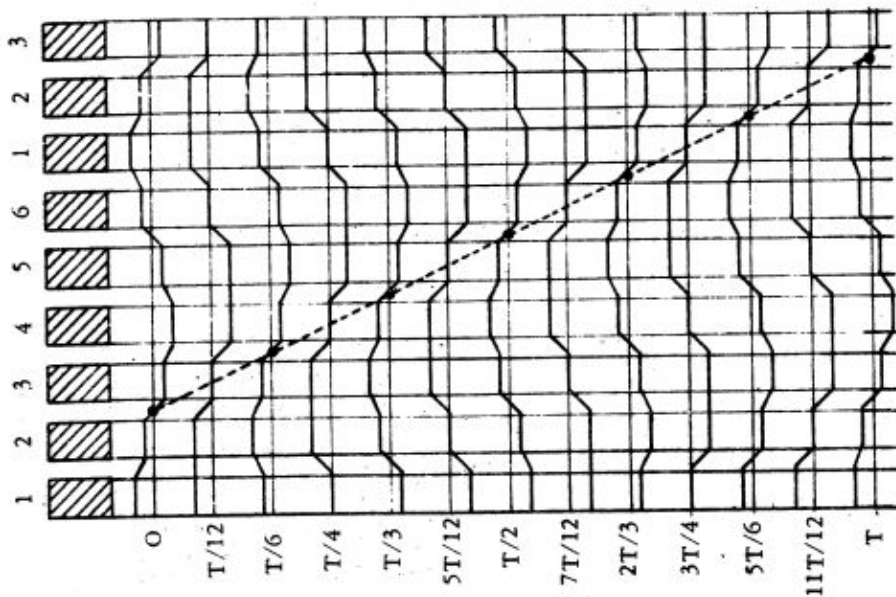


Fig. 9-2

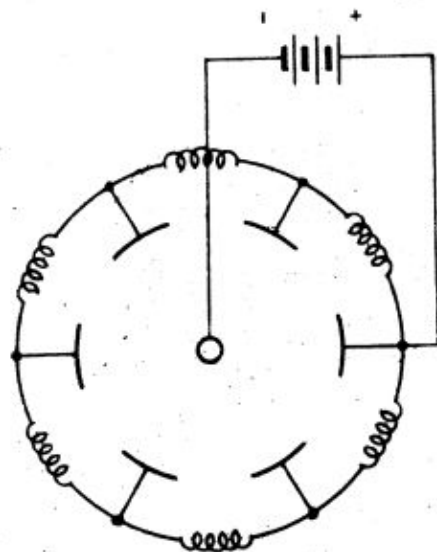


Fig. 9-1.

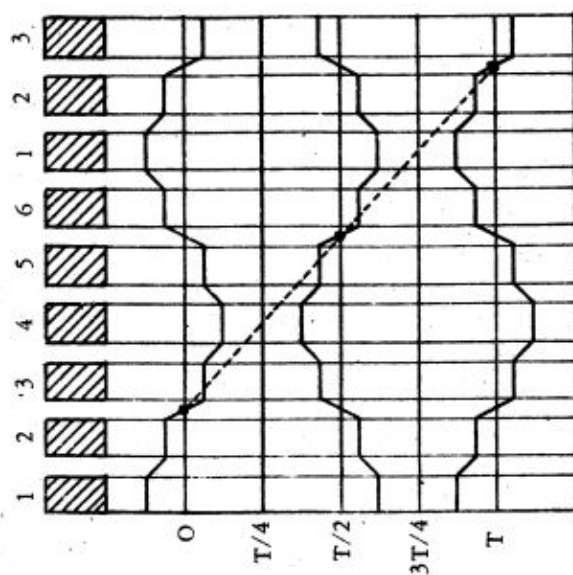


Fig. 9-3

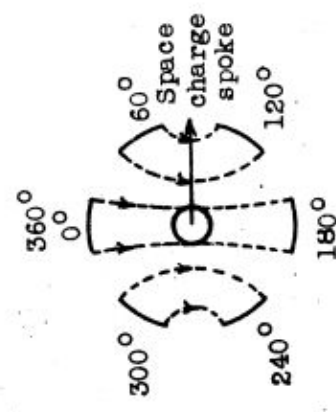


Fig. 9-4

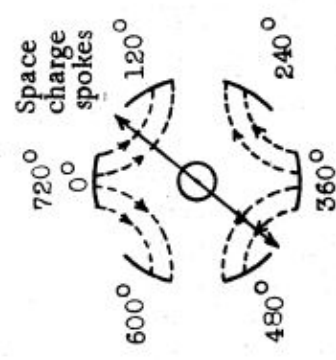


Fig. 9-5

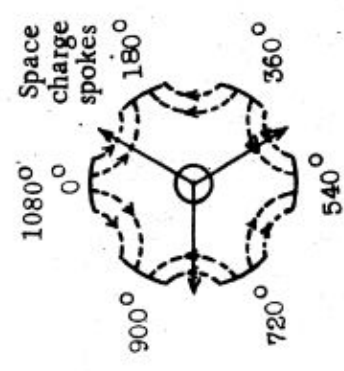


Fig. 9-6

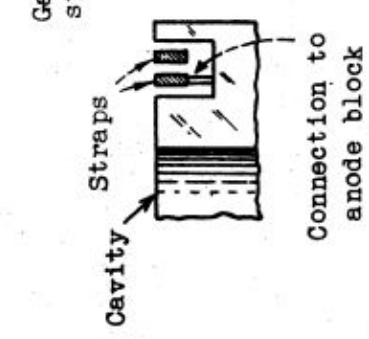


Fig. 9-7

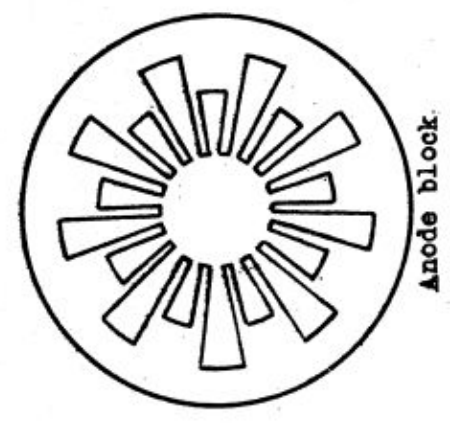


Fig. 9-8

RADIO NAVIGATIONAL AIDS.

symmetrical about a vertical centre line. The maximum field across a gap occurs at the sides and at the right the field is retarding. The space charge wheel has therefore a single spoke, in the position shown by the arrow; it makes one revolution per cycle.

The second mode is shown in Fig. 9-5. The field divides into four parts alternately accelerating and retarding, so that the space charge wheel has two spokes in the positions of maximum retarding field shown, and makes one half revolution per cycle.

The third mode is shown in Fig. 9-6 and has six sections. The space charge wheel then has all three spokes, spaced 120° apart, and rotates 120° in each cycle.

The fields shown are considerably simplified, containing only the fundamental components, so that they represent conditions not too close to the anode surface.

In addition there is a mode, called the zero mode, in which all anode segments are in phase. The anode potential has the same value everywhere, and therefore the r-f field has no horizontal component at the anode surface. The magnetron then behaves as if it had a complete cylinder without slots as anode, and can only operate at the cyclotron frequency. To avoid this mode, the magnetic field is so adjusted that the cyclotron frequency is considerably different from the frequency of the coupled resonator system.

The most important mode is that in which successive segments are opposite in phase, or the phase difference is 180° or π radians. With different numbers of anode segments, the numbering of this mode corresponds to half the number of segments, so that with six segments it is the third mode. It is therefore convenient to call it the π mode, as this term can be used with any number of segments.

9-3 Importance of Pi Mode.

For a given peak value of r-f field, the greatest difference of potential across the gaps occurs for the π mode, since there is a phase change of 180° between segments, amounting to a complete reversal. Thus the resonators will be excited most strongly, and all equally. The desirable condition is that as soon as the H. T. is applied, oscillation should take place in the π mode and no other.

Each mode can only be driven by the electrons at certain values of H. T. and magnetic field. But if the various mode frequencies are grouped very closely it is possible for forced oscillations to be produced, at the frequency of the driven mode, but actually in a field pattern corresponding to another mode. The driven mode is then said to be "contaminated" by the other mode. Modes that are not contaminated are said to be "clean".

In order to obtain a clean oscillation in the π mode, it is desirable to separate the mode frequencies as far as possible. For this purpose, use is made of a method known as "strapping".

9-4 STRAPPING.

In the π mode, the anode segments are alternately positive and negative as regards rf. The positive segments are all at one potential, the negative segments at another. Thus the positive segments may all be connected together by conducting "straps" at the top and bottom, and the negative segments may be connected by a second set of straps.

In other modes the alternate segments are not at the same potential. Thus in the first mode for a six-segment anode, alternate segments differ in phase by 120° . Connecting these segments together tends to stop such a mode from being produced.

The zero mode is not affected by strapping. If all segments are at the same potential, it is of no assistance to divide them into two strapped sets.

The straps are not in the form of complete circles, but have gaps carefully located with respect to the position of the output coupling loop. All modes except the zero and π modes are "double", which means that the phases may lag either in one direction or the other. The coupling loop tends to produce reflections and therefore standing wave systems, and the strapping systems used are so arranged as to reduce equally both components of the double modes. The theory here is beyond the scope of the course.

Several systems of strapping have been used. A typical arrangement for an eight-segment anode is shown in Fig. 9-7. The top view is shown at the right. There are two rings, each containing a gap, and these are connected to alternate segments. The gaps are one on each side of the resonator opposite to that containing the coupling loop. The method of connection is shown at the left. The straps are placed in a channel at the end of the anode block, near the inner edge, and are connected to the segments as required by means of pins. The same arrangement may be repeated at the bottom. Instead of strapping we may use an anode block of the form shown in Fig. 9-8, known as the "rising sun" pattern. Two shapes of resonator are used, long and short, alternately. This idea amounts to staggering the resonant frequencies to produce an overall double-peaked resonance curve, instead of tightening the coupling.

The rising-sun arrangement favours the π mode, in which all the long resonators are in one phase and all the short ones in opposite phase, in comparison with lower modes, in which associated resonators must be selected from both groups. Strapping is unnecessary. This pattern is particularly adapted to very high frequencies, as a large number of resonators can be used, 18 in the case shown. The zero mode is not suppressed, and usually gives a noticeable contamination of the π mode. This can only be avoided by choice of operating conditions, especially the value of the magnetic field, so that the cyclotron frequency differs from the natural frequency of the resonators.

9-5 Tunable Magnetrons.

A magnetron is essentially a fixed-frequency oscillator, but it is possible to provide for tuning over a limited range. Fig. 9-9 (a) shows a method of

RADIO NAVIGATIONAL AIDS.

tuning by inserting brass slugs in one end of the resonators. This has the effect of reducing the inductance of the equivalent tuned circuit. Fig. 9-9 (b) shows a method by which the slot capacitance can be varied, and Fig. 9-9 (c) shows variation of the capacitance between straps. In each case the tuning mechanism is controlled by turning a knob located outside the vacuum. A variation of from 5 to 10 percent from the mid-frequency is obtainable by any of these methods. Tunable magnetrons are especially useful for replacements, as it is easier to tune the magnetron than the remainder of the set.

9-2 Output Circuit.

The commonest method of taking output from a magnetron is by means of a loop coupled into one of the resonators as shown in Fig. 9-7. This loop is the end of the central conductor of a coaxial line, and must be arranged to match the line to the resonator system.

In practice the resonators are so small that the coaxial line would have too small a diameter. Fig. 9-10 shows an arrangement that permits coupling to a larger line. It contains a short section of two coaxial cones. At any cross-section the ratio of the radii is constant, and so the characteristic resistance is uniform so that no reflection is produced. This conical section is inside the vacuum, and so a higher voltage can be allowed between the cones without breakdown than would be possible in air. The greatest diameters of the cones are made suitable for connection to a coaxial cylindrical line outside the vacuum.

The "choke" connection to the output line contains folded low-impedance sections built into each conductor. In the outer conductor a half-wave line folded once is used. This is short-circuited at the end and so reflects a short-circuit across its mouth. In the inner conductor there is a quarter-wave folded section terminating in an open circuit at the end of the inner post, again reflecting a short-circuit to its mouth. In neither conductor is there any reliance on metal-to-metal contact.

For very high frequencies, the output may be taken by means of a wave guide instead of a loop, a quarter-wave transformer section being used for matching, as shown in Fig. 9-11. This section is built into the copper resonator block.

9-7 MAGNETRON CHARACTERISTICS.

The most important characteristic of a magnetron is its volt-ampere characteristic, shown in Fig. 9-12. The d-c current is used as abscissa and the d-c voltage as ordinate. On the chart are drawn lines of constant power output (dot-dash lines, the numbers representing kilowatts), constant magnetic field (solid lines, the numbers representing gauss), and constant efficiency (dash line, percentages).

The magnetic field and the d-c voltage are the independent variables, and the current is determined by them. To obtain more current at a given magnetic field, the voltage must be increased, and so the lines of constant magnetic field slope upwards, they are practically straight and parallel.

The lines of constant power output are curved, and slope downwards to the

right. The curve for lowest power output has a minimum ordinate in the range of the diagram, and the voltage rises again for higher currents. This is because the efficiency is changing. The minimum ordinates for the other curves fall outside the operating range.

The power input is the product of current and voltage, and so the efficiency can be calculated at any point. Each efficiency curve shows a well-defined minimum ordinate. For a fixed voltage, the efficiency at first rises as the current increases, then reaches a maximum and falls again. This can be seen clearly at the 20-kilovolt level. The combination of a relatively strong magnetic field, high supply voltage, and moderate current gives good output at good efficiency. The maximum efficiency varies from about 50 percent to over 80 percent in different types of magnetrons.

For comparison, Fig. 9-13 shows the characteristics for a rising-sun magnetron in which the π mode is contaminated by the zero mode. The field strength 3920 gauss corresponds to the cyclotron frequency, and the zero mode is then strongly excited, giving lowered efficiency. High efficiency can be obtained either with higher voltage and field, or with a relatively low voltage and a weaker field. The power output curves show a marked bulge to the right in the region of low efficiency. In the lower part of the diagram the efficiency curves are totally different in shape from those in Fig. 9-12.

9-8 Load and Frequency.

In a strapped magnetron, in which all modes other than the desired π mode are suppressed, the resonator system behaves approximately as a single tuned circuit.

A tuned circuit has two quantities associated with it, namely the resonant frequency and the Q factor, which is a measure of the sharpness of resonance. Q is equal to 2π times the ratio of stored energy loss per cycle, and for an unloaded resonator is usually well over 10,000.

When a load is coupled to the resonator, the losses are increased and the effective value of Q falls. We may therefore define three separate Q factors - Q_0 for the unloaded resonator, Q_L for the resonator with the load connected, and Q_{ext} , the "external" Q, which is 2π times the ratio of energy stored in the resonator to that dissipated in the load per cycle. The relation between the three values of Q is -

$$\frac{1}{Q_L} = \frac{1}{Q_0} + \frac{1}{Q_{ext}}$$

The ratio of Q_L to Q_{ext} is equal to the ratio of power dissipated in the load to that dissipated in the load and resonator together, and is known as the circuit efficiency. To obtain the overall efficiency, this must be multiplied by the electronic efficiency of the magnetron, which is the efficiency of converting dc into rf. Variations of load affect mainly the circuit efficiency.

The actual load is connected to the coupling loop by a coaxial line or wave guide. Even if it is a pure resistance, unless it matches the line exactly, it may

RADIO NAVIGATIONAL AIDS

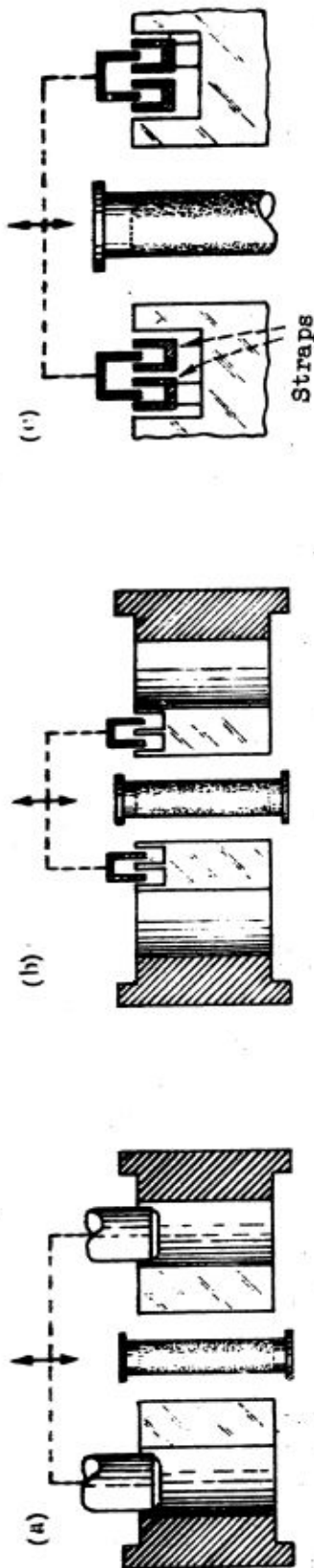


Fig. 9-9

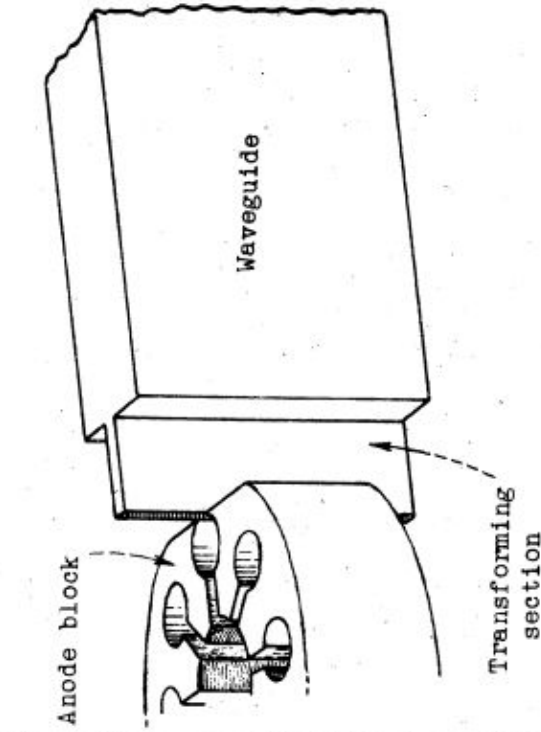


Fig. 9-10.

Fig. 9-11

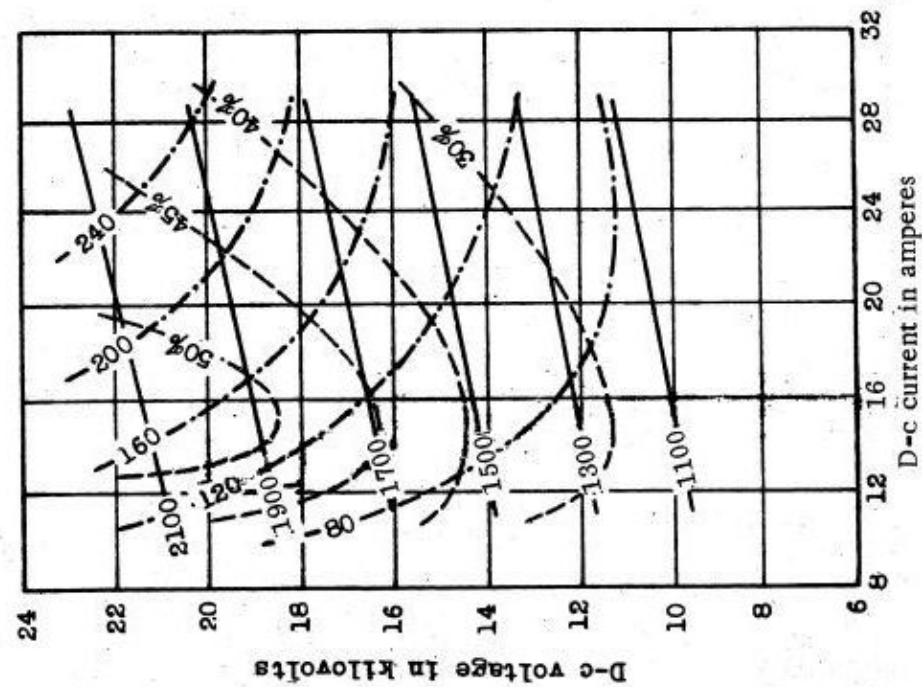


Fig. 9-12

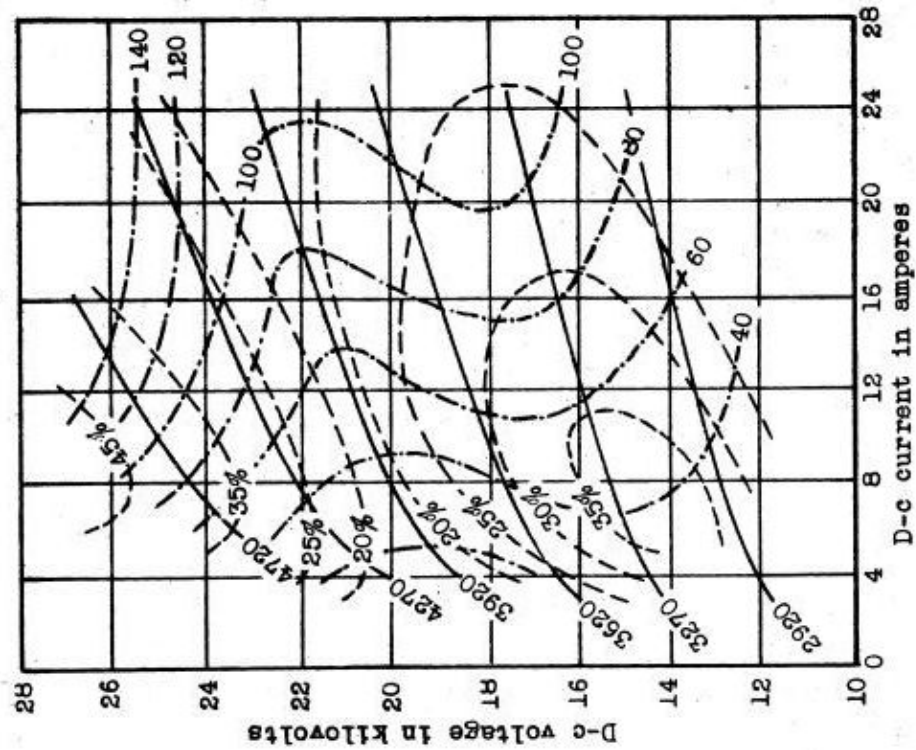


Fig. 9-13

RADIO NAVIGATIONAL AIDS.

not appear as a resistance at the coupling loop, but as an impedance varying with the length of the guide. Any reactive component will have an effect on the frequency of resonance, because reactance will be reflected into the resonator.

With perfect matching the guide length makes no difference, but this condition is difficult to obtain. It is common to assume a reflection coefficient at the load of 20 percent, giving a standing wave ratio of 1.5 to 1. For such a load, the frequency will vary between a maximum and minimum value as the guide length is changed. The difference between the maximum and minimum frequency for 1.5 to 1 S.W.R. is known as the "pulling figure" and is a convenient single figure to express the frequency stability.

9-9 STARTING OF OSCILLATIONS.

The magnetron is required to commence oscillation during every HT pulse, and to oscillate correctly in the desired mode and no other.

The oscillations in the resonators take some time to build up, especially when the value of QL (with the load connected) is high. Oscillation in any particular mode will only take place over a narrow range of values of H. T. when the magnetic field is fixed. It follows that the leading edge of the applied pulse should not be too steep, as otherwise the voltage may grow more rapidly than the oscillations can be built up.

In many cases the power supply (modulator) has poor regulation, and there is a substantial voltage drop in the modulator under normal operating conditions, depending on the magnetron current. This current can only flow when oscillation takes place, and so if oscillations fail to start the voltage rises above normal, and often beyond the range of voltages that give oscillation.

If, then, the leading edge of the pulse is too steep, the voltage may become too high for oscillation before the oscillation can build up. In some cases however, an oscillation will then start in a different mode, requiring higher voltage.

Another possibility is a momentary oscillation at a lower voltage in an undesired mode during the application of the leading edge, followed by correct oscillation in the desired mode when normal voltage is reached. This type of incorrect moding does not cause much operating trouble, and does less damage than the other, which causes additional strain on the insulation.

9-10 MAGNETIC CIRCUIT.

For experimental work and the measurement of characteristics, electromagnets must be used. When the best value of operating field has been found it is more convenient to use a permanent magnet. For the longer waves a separate magnet is used, but with shorter waves a lighter structure is obtained by building the magnet into the magnetron structure itself. The required magnetic fields are very strong, but extend over only a small area.

QUESTION PAPER FOR CHAPTER 9.

1. Explain what is meant by strapping in a magnetron, and what is its purpose.
2. What is the π mode, and why is it important ?
3. Describe two methods of coupling the output circuit to a magnetron.
4. Describe the characteristic curves of a magnetron. What are the symptoms of zero moding ?
5. Describe three methods by which magnetrons can be made tunable to different frequencies.

MARCONI SCHOOL OF WIRELESS.
RADIO NAVIGATIONAL AIDS.
CHAPTER 10 - PULSE SHAPING.

10-1 INTRODUCTION.

A large number of the circuits in radar equipment are concerned with the shaping of pulses to produce the desired waveform. It is frequently not possible to produce with one circuit the exact waveform required. It is therefore necessary to use auxiliary circuits which, after the approximate waveform has been produced, are capable of reforming it. In this way a pulse may travel through several different stages of shaping before the desired result is obtained.

For convenience in study, shaping circuits can be roughly divided into two main types (1) those in which the shaping depends on valve characteristics, as for example when diodes and overdriven amplifiers are used to remove an unwanted portion of a waveform. An example of this type of shaping is shown in Fig. 10-1 (a) being the pulse as originally passed into a shaping circuit, and (b) the pulse after shaping; (2) those in which the shaping depends on combination of circuit elements, containing R and C elements, L and R elements, or a combination of all three. An example of this type of shaping is shown in Fig. 10-2, (a) being the pulse as originally passed into a shaping circuit, and (b) the pulse after shaping.

The main operations performed by shaping circuits are of three types, known as "clipping", "differentiation", and "integration". Clipping tends to produce a square wave from an input wave of any shape. Differentiating circuits or "peakers" tend to produce short pulses at instants when the input voltage is changing rapidly. Integrating circuits are the reverse of differentiating circuits, and with a square wave input tend to produce a triangular wave, with a constant rate of increase or decrease.

10-2 Clipping with Diode.

Clipping may be obtained by the use of biased diodes, which operate on one half-cycle only. Fig. 10-3 (a) shows a suitable circuit diagram. The input voltage is in series with a resistance that is large in comparison with the resistance of the diode when conducting, and the series circuit contains a d-c voltage that determines the level at which clipping occurs.

Whenever the input voltage is negative, or positive and less than the d-c

bias, the diode is non-conducting, and the output voltage is practically equal to the input voltage. When the input voltage is positive and greater than the bias, diode conducts and may be regarded as a short-circuit which fixes the output voltage at the battery emf. Thus a flat top is produced in the output wave, as shown in Fig. 10-3, (b). Two diodes may be used in succession, operating on opposite half-cycles. The top and bottom will then both be flattened, but the intermediate part will be unchanged in shape, as in Fig. 10-3 (c). However a truly square wave will not be obtained.

10-3 Clipping with Triode or Pentode.

A better result may be obtained by using a triode or pentode amplifier, which clips both top and bottom at the same time. The valve is, by ordinary standards, severely overdriven; up to 100 volts input signal may be applied to an ordinary receiving valve. This signal is, however, applied through a high resistance, for example 100,000 ohms. The circuit is shown in Fig. 10-4, and is standard except for the series grid resistor. The valve is operated at zero bias.

To determine the performance of this circuit, we must first find the actual waveform of the voltage at the grid. On the positive half-cycle, the grid collects current, as shown by the $E_g - I_g$ static characteristics of Fig. 10-5. Below is shown the positive half-cycle of the input wave E_i . At the peak is shown a construction to find the actual grid potential. The input voltage is set off along the zero current line, and a load line is drawn from this point at a slope corresponding to the series resistor R_s . The intersection with the grid current curve has an abscissa the actual peak grid voltage, which is seen to be very small. Thus the positive half-cycle of the input voltage is almost completely suppressed, while the negative half-cycle is unaffected.

The plate current may now be determined, by the ordinary construction shown in Fig. 10-6. It is seen that the plate current is cut off for nearly the whole of the negative half-cycle, and the curve rises and falls very steeply. The top is slightly curved, as the grid goes somewhat positive.

In the output voltage, which is opposite in phase to the plate current, the top is flat and the bottom slightly curved. If a second similar stage follows, both top and bottom will be flat and the sides practically vertical, giving a rectangular wave.

If bias is applied to the grid in Fig. 10-4, the peak grid voltage will still be limited to approximately zero, and the main effect will be that the conducting and non-conducting periods will have different lengths. The relative lengths can readily be controlled over a certain range by the bias, as in Fig. 10-7. With positive bias the conducting period is lengthened, with negative bias shortened.

10-4 Resistance-capacitance Circuits.

Resistors and valves alone can amplify an input wave without distortion, or can shape it to be nearly rectangular. To obtain other waveforms it is necessary to make use of reactive elements. The most important combination consists of a resistance R and a condenser C . Suppose these to be connected in

RADIO NAVIGATIONAL AIDS

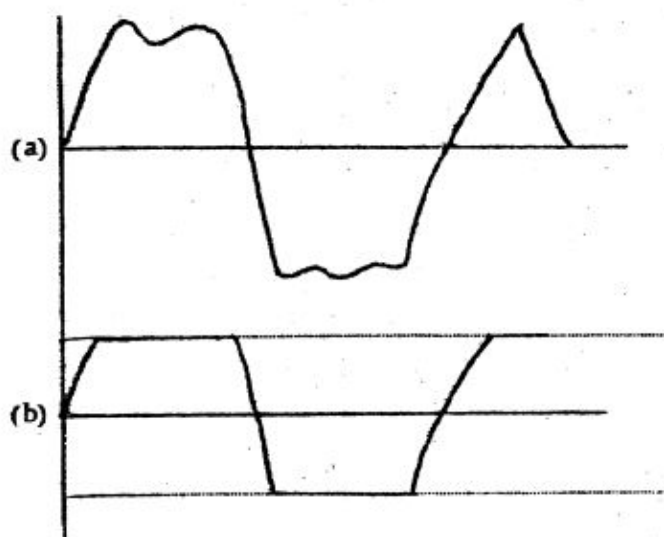


Fig. 10-1

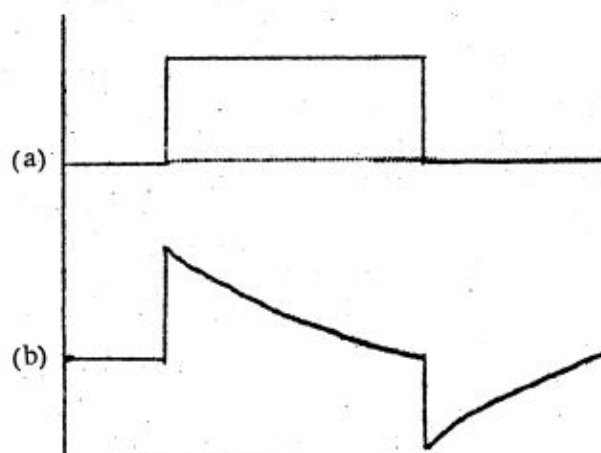


Fig. 10-2

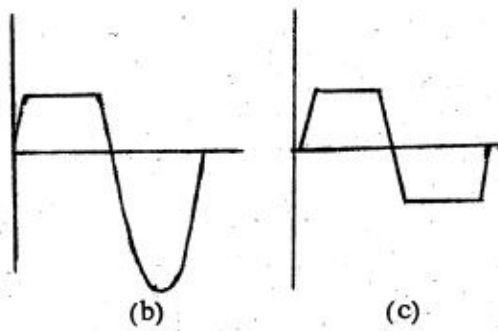
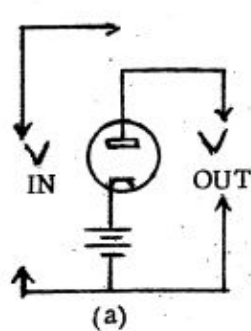


Fig. 10-3

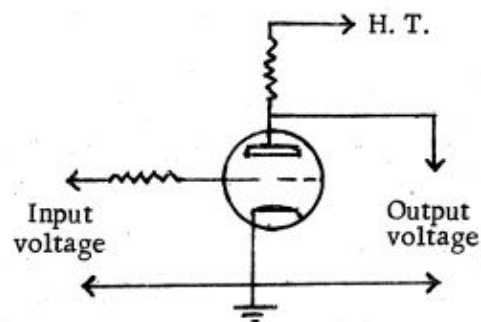


Fig. 10-4

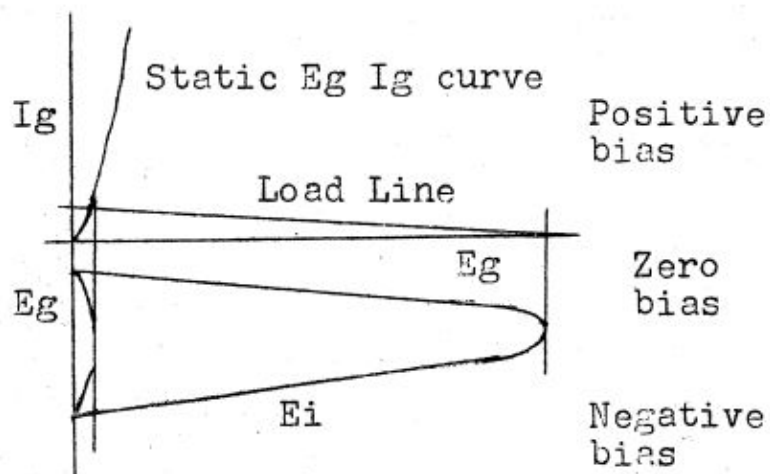


Fig. 10-5

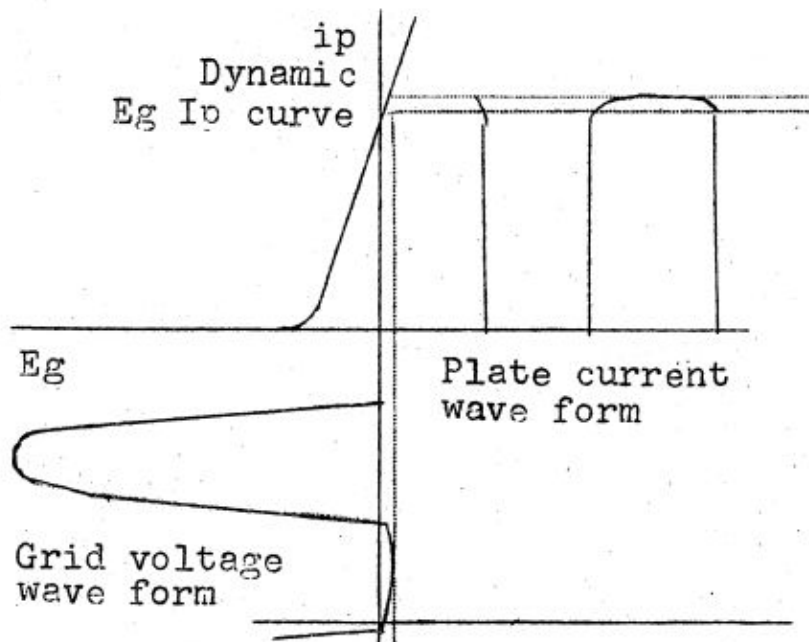


Fig. 10-6

RADIO NAVIGATIONAL AIDS

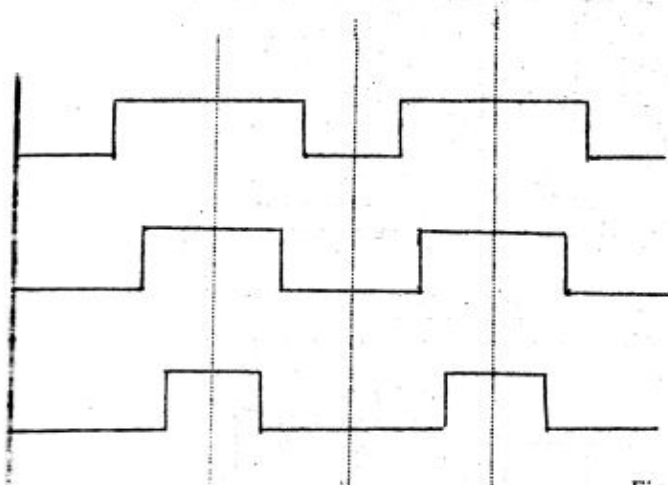


Plate current wave forms

Fig. 10-7

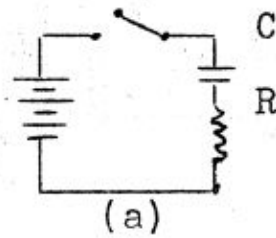


Fig. 10-8

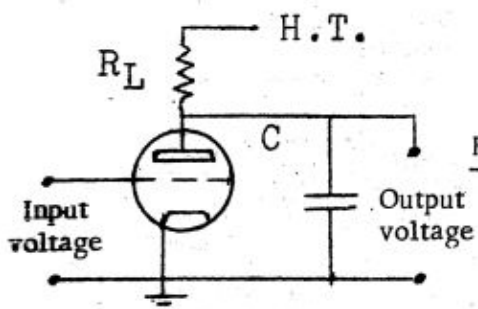
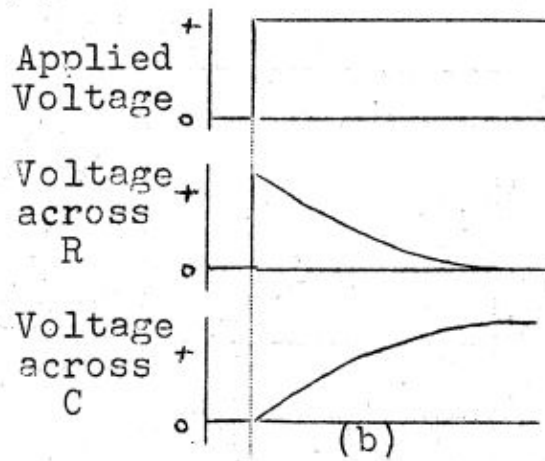


Fig. 10-9

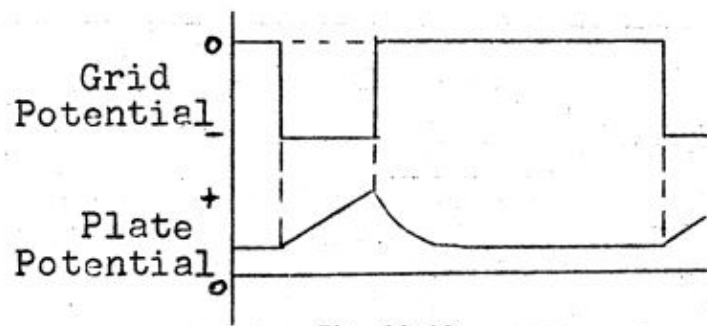


Fig. 10-10

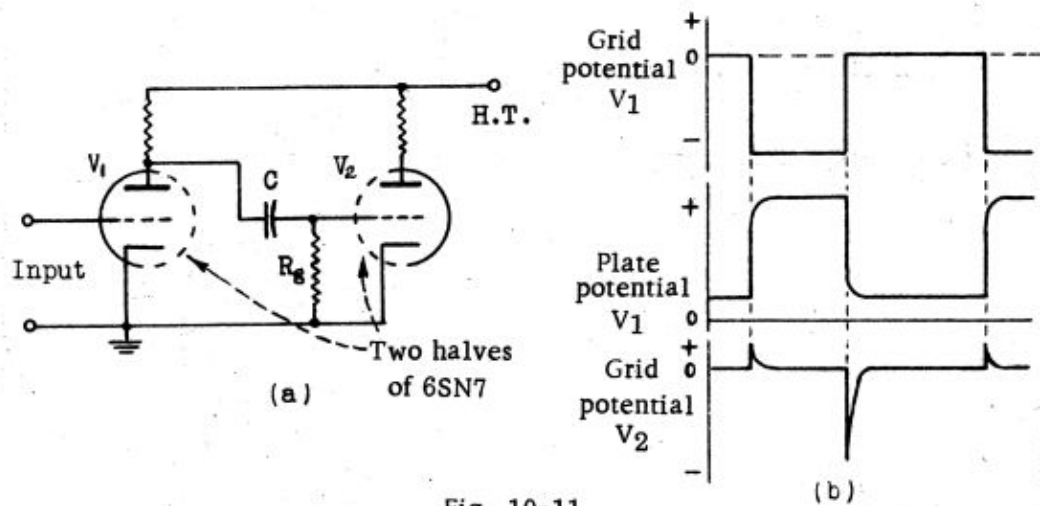


Fig. 10-11

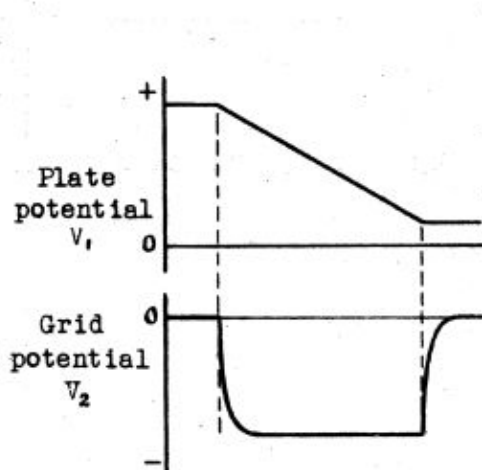


Fig. 10-12

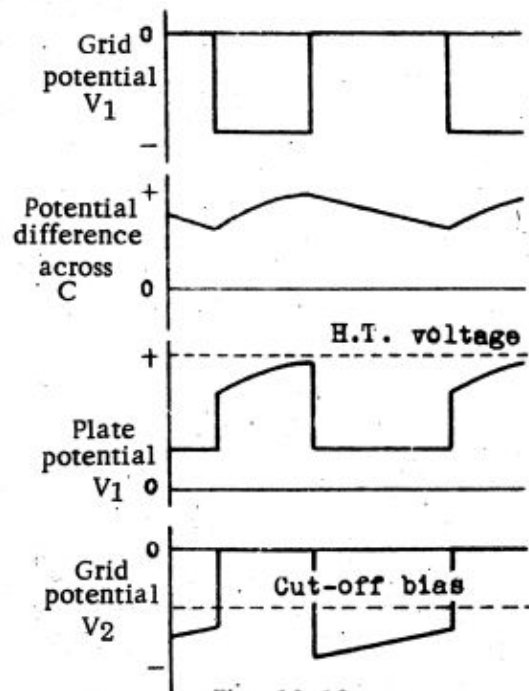


Fig. 10-13

RADIO NAVIGATIONAL AIDS.

series with a switch and a d-c voltage, as in Fig. 10-8 (a), than the total and individual voltages obtained on closing the switch are as shown in Fig. 10-8 (b). The condenser is uncharged at the beginning, and cannot suddenly change in voltage, so that the whole of the applied voltage is impressed across the resistance R , and the initial current is maximum. This current charges the condenser, and the voltage built up across the condenser reduces that across the resistance. The current therefore gradually falls. This means that the condenser is charging less rapidly.

The curve of voltage drop across the resistance, which is of the same shape as the current curve, is called an exponential curve. It has the property that the change in any given interval of time is proportional to the value at the beginning of that interval. Thus if the curve falls to half the initial value in 10 μ s, it will fall to half this or a quarter of the initial value after a further interval of 10 μ s, making 20 μ s from the beginning.

The rate at which the current starts to fall can be calculated. The initial current is $\frac{E}{R}$ amperes, flowing into a capacitance of C farads. This builds up charge at a rate of $\frac{E}{R}$ coulombs per second, so that the condenser voltage rises at a rate of $\frac{E}{CR}$ volts per second, and the voltage across the resistor falls at the same rate. At any later stage, the rate of fall of voltage is reduced to $\frac{E^I}{CR}$ volts per second, where E^I is the instantaneous voltage; this can be seen by the same argument.

If the initial rate of fall of the voltage across the resistor could be maintained, the voltage would fall to zero after a time CR seconds. This is known as the "time constant" of the circuit. Actually, after an interval equal to the time constant, a fraction equal to .37 of the original voltage will remain. After a further interval equal to the time constant, the fraction remaining is reduced to $.37 \times .37 = .14$. If we call T the time constant, then after $3T$ we have .05 left, and after $5T$ less than .007 left.

Suppose next that the switch is kept closed, and the applied voltage is suddenly changed to a different value. The resistance will take the whole of this change at the beginning, but as the condenser charges the whole of the change is transferred to it, and the current falls. The change will be complete within less than 1 percent after an interval $5T$; the same fractions apply as before.

10-5 Production of Saw-Tooth Wave-Form.

Suppose that a square wave is applied in a negative direction to the grid of the valve shown in Fig. 10-9, which has as load a resistor R_L in series with the H. T. supply, and a by-pass condenser C between the plate and the cathode, which is earthed. The input wave alternately biases the grid beyond cut-off for one interval and brings in to a zero potential for another interval.

When the grid has been at zero potential for a considerable time, the plate voltage is determined by the valve characteristics, the load resistance, and the supply. If the load resistance is high, the plate voltage will be low. Now the condenser voltage is equal to the plate voltage, owing to the circuit arrangement and consequently the condenser is nearly discharged.

When the valve is cut off, the condenser is effectively in series with the load resistance across the H. T. supply, and therefore begins to charge up towards the full supply voltage, according to an exponential curve with time constant CR_L seconds. It is not allowed sufficient time to flatten out the curve, but the valve is suddenly brought to a conducting state again. The effective resistance across the condenser consists of R_L and R_p in parallel, and is a little lower than R_p . (R_p is between plate and earth, R_L between plate and H. T. which is at earth potential as regards varying components). If this resistance is denoted by R , the condenser discharges with time constant CR seconds. The discharge curve is allowed to flatten out completely, as in Fig. 10-10, to a final voltage the same as the initial voltage before charging commenced. This completes one cycle of the sawtooth wave-form; one cycle is produced for each negative pulse on the grid.

The value of R_p to be used is considerably higher than that for normal amplification. Owing to the high load resistance, the load line cuts the characteristic curves of the valve well down on the lower bend, and R_p may readily be twice the normal value.

If the time during which the condenser is allowed to charge is small in comparison with the time constant for charging, the rising part of the saw-tooth will be nearly straight. Its duration is the same as that of the applied negative grid pulse. The falling part of the saw-tooth is curved, but this is not of any importance and no attempt is made to straighten it.

The above circuit belongs to the class of integrating circuits, as defined earlier. Differentiating circuits can also be made with resistance-capacitance arrangements.

10-6 Resistance-capacitance Peaker.

Narrow pulses may be obtained from rectangular waves by the use of a resistance-capacitance coupled amplifier in which the coupling condenser and grid leak are made very low, so as to give a low time-constant. A twin valve such as type 6SN7, equivalent to two 6J5 valves in one envelope, is convenient. The circuit is arranged as in Fig. 10-11 (a).

The grid is again pulsed between zero potential and a large bias beyond cut-off, the wave-form being as shown in the upper diagram of Fig. 10-11 (b). When the valve V1 has been conducting for a long time the plate voltage is low. As a steady state has been reached, there is no current in the grid leak and so the grid of V2 is at earth potential or zero bias, and the valve is in the conducting state.

When a negative pulse is applied to the grid of V1, the plate current is immediately cut off. The load across the supply then consists of the load resistance, the condenser, and the grid leak. The difference between the supply voltage and the condenser voltage must appear as voltage drops across the two resistors, in proportion to their resistances. Thus the right hand plate of the condenser is no longer at earth potential, but jumps suddenly to a positive value, the other plate jumping by the same amount.

RADIO NAVIGATIONAL AIDS.

As a result the grid of V_2 becomes positive, and grid current flows. This amounts to a shunting of the grid leak, and so the grid of V_2 does not jump as much as it would without grid current. At the same time the potential of V_1 jumps by an equal amount, since the condenser voltage is unchanged. The condenser now charges, and raises the plate potential of V_1 gradually. When the process is complete and the current has ceased, the plate potential of V_1 has risen to the supply voltage, as in the middle diagram of Fig. 10-11 (b) while the grid potential of V_2 has returned to zero, as in the lowest diagram. The time constant is obtained as the product of C and the effective resistance

$$R_L + \frac{R_g r_g}{R_g + r_g}, \text{ where } r_g \text{ is the internal grid-cathode resistance of } V_2$$

with positive grid. If C is small enough, the time constant may be a few microseconds only.

When the pulse ends, the grid of V_1 is brought back sharply to zero potential, and V_1 again conducts. The plate potential of V_1 is suddenly lowered and as the condenser voltage cannot change immediately the grid of V_2 is driven strongly negative, and V_2 is cut off. After a few microseconds this grid returns to its normal condition of zero potential. The time constant is obtained as the product of C and a resistance

$$R_g + \frac{R_L R_p}{R_L + R_p} \text{ where } R_p \text{ is the d-c}$$

plate resistance of V_1 , which here effectively shunts the load resistance R_L as regards condenser discharge. On the other hand R_g is unshunted, since the grid current of V_2 is cut off.

The plate voltage of the first valve shows an amplified copy of the input pulses, but with rounded corners. The amount of rounding depends on the capacitance of the condenser, or, more exactly, on the time constants, which are different for the two corners affected. These time constants must be kept small for good pulse shape.

The grid voltage of the second valve has narrow pulses, which occur whenever there is a sudden jump in the input wave. The response is not to the actual input voltage, but to changes of input voltage.

When this circuit is used with a sloping wave-form, instead of a square wave-form, the results depend on the relative values of time constant and duration of the slope. If the slope is long in comparison with the time constant, the result is a nearly square pulse, with rounded corners, as in Fig. 10-12. The magnitude of the square pulse is then proportional to the slope of the inclined portion of the input wave. A peaking circuit may thus be used to convert saw-tooth waves into rectangular pulses.

10-7 Resistance-capacitance Amplifier.

If the coupling condenser is made large enough, the time constants become longer than the pulse duration. The values used for coupling condenser and grid leak are such as would be suitable for a a-f sine wave amplification. The circuit arrangement is the same as in Fig. 10-11 (a). During the period when

V_1 is non-conducting, the condenser C is charging, but when V_1 conducts the condenser is able to discharge. During charge it tends to rise to the full H. T. voltage as a limiting value, but during discharge it tends to fall to the plate voltage of V_1 in the conducting state as limiting value. The time constants for charge and discharge are too long for these processes to become complete, and a state of balance is reached in which the charge and discharge are equal in any cycle, as shown in Fig. 10-13.

When V_1 is non-conducting, the grid of V_2 is nearly at zero bias. The plate voltage of V_1 is therefore practically the same as the voltage across the condenser. But during the conducting period, the plate voltage of V_1 has a practically constant lower value. Thus the plate voltage of V_1 suddenly jumps down at the beginning of the conducting period, and suddenly jumps up at the end. The grid voltage of V_2 has jumps of the same amount. It cannot rise appreciably above zero, so the top of the wave-form of this quantity is practically flat at zero, while the bottom slopes upwards. The whole of this slope may be made to lie below cut-off, and the plate current of V_2 (not shown) is practically an undistorted square wave. The output may be taken between the plate of V_2 and earth.

The circuit of Fig. 10-11 (a) is thus a peaker or an amplifier, according to the time constants involved.

During the time when V_2 is conducting its grid is at zero potential. The average grid voltage taken over a cycle is negative, and thus the grid is effectively biased. This type of bias is sometimes called signal bias, but is simply grid-leak bias, similar to that occurring in ordinary Class C amplifiers. A voltage across the grid leak may occur from pulses through the condenser or through flow of grid current. Over any number of complete cycles the coupling condenser returns to the same state of charge, and so the average current flow in one direction through the condenser is the same as that in the opposite direction. Thus the average current through the grid leak does not come from the condenser, but is grid current. The grid rises above zero potential when V_2 conducts, by a sufficient amount to collect this grid current, a few volts being adequate.

In the circuit of Fig. 10-11 (a) the grid leak of V_2 is returned to the cathode, so that the only bias is due to grid current. It is also possible to apply bias by returning the grid leak to a point at negative potential. In this way grid current can be prevented entirely if the grid is not driven positive. The waveform of the grid voltage of V_2 is then of the form shown in Fig. 10-14. It has the disadvantage that the top of the wave is decreasing towards the bias as a limiting value, and the output wave in the plate circuit will have a similar slope on the top.

Frequently the grid return is taken to +H. T. through a very high resistance, as in Fig. 10-15 (a). Suppose the supply to be 300 v. and the grid leak 1 megohm. Then if there is no input signal, the current in the grid leak is about 300 μ a or .3 ma., and the grid is only very slightly positive. In operation the grid is driven only slightly positive at the commencement of the top of the wave, and

RADIO NAVIGATIONAL AIDS

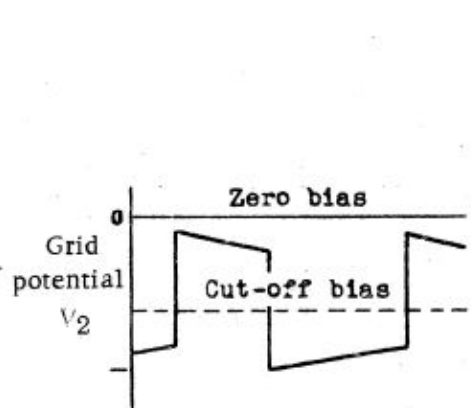


Fig. 10-14

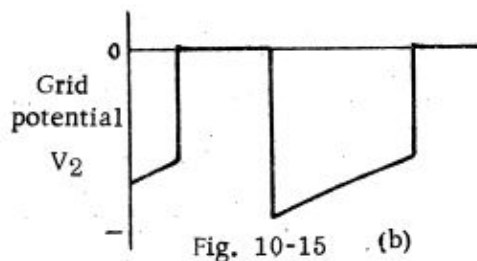
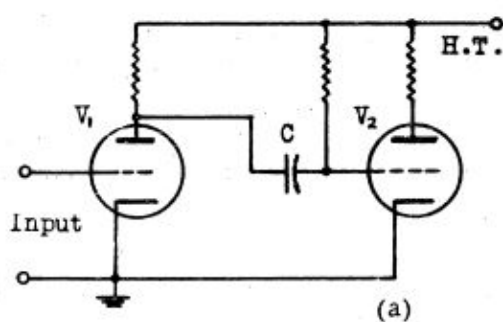


Fig. 10-15 (b)

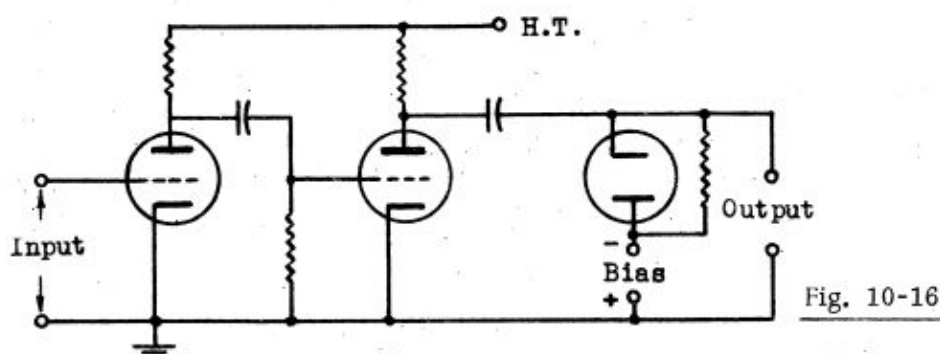


Fig. 10-16

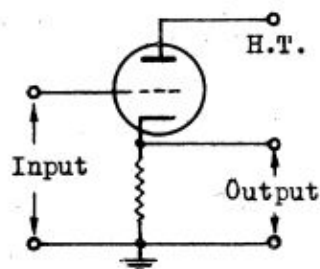
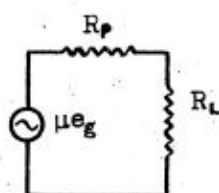
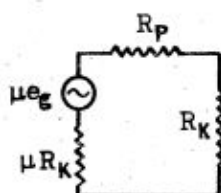


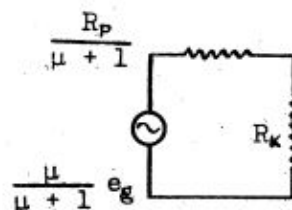
Fig. 10-17



(a)



(b)



(c)

Fig. 10-18

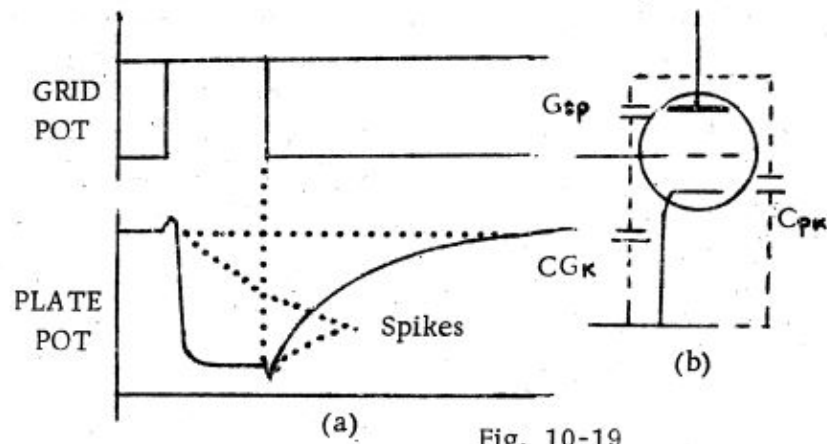


Fig. 10-19

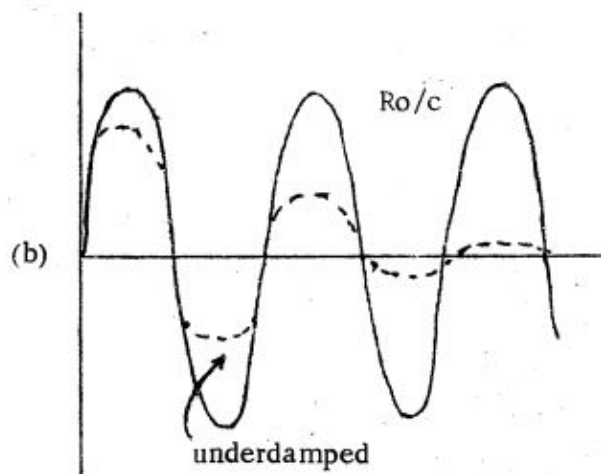
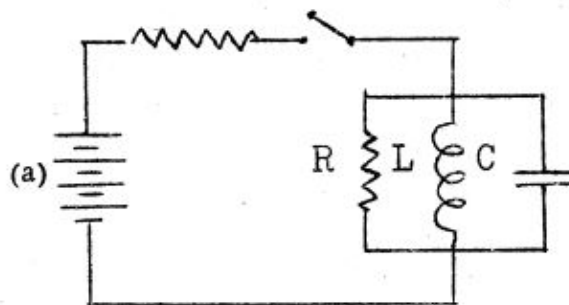
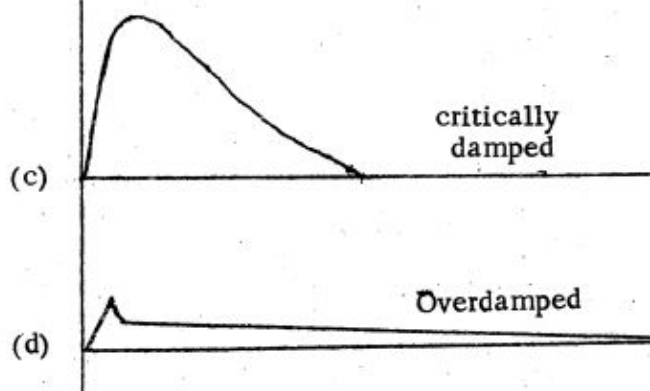


Fig. 10-20



RADIO NAVIGATIONAL AIDS.

remains substantially at zero volts during the whole of this portion. The slope of the bottom part of the wave is increased, since the grid is returning towards +H. T. as limiting value instead of towards zero voltage. This change is, however, partly counteracted by the increase of time constant due to the higher grid-leak resistance. The wave-form of the grid voltage is indicated in Fig. 10-15 (b).

10-8 Clamping.

Whenever a varying voltage is transferred from one circuit to another, through a condenser or a transformer, only the varying component is transferred, and not the d-c component. The output wave can be moved bodily up or down by applying a d-c voltage.

In amplifiers, the bias is so arranged that the wave-form can be amplified without clipping through cut-off or grid current. In the output circuit, however, it is often necessary to "restore" a lost d-c component. This is done by means of a "clamping" circuit, which fixes the potential of either the top or bottom of the wave.

For example, the sawtooth used in a sweep circuit for a cathode-ray tube is clamped to begin always at the same voltage, corresponding to a fixed abscissa on the cathode-ray tube screen. This is necessary to make sure that successive traces will coincide.

In the curves of Fig. 10-13, the top of the grid voltage wave of V_2 is approximately clamped to zero voltage, as grid current flows when the grid becomes positive.

The circuit of Fig. 10-16 shows how a diode may be used to clamp the bottom of a wave-form to any voltage. The diode is connected with its cathode to the positive output lead, and with its plate to a point at the voltage to which the bottom of the wave is to be clamped. The cathode lead may become positive without hindrance, but if it becomes more negative than the plate, the diode will conduct. The current flows to charge the coupling condenser and raise the cathode potential, so that the cathode is effectively prevented from becoming more negative than the clamping voltage. It is necessary to have a high resistance in parallel with the diode, so that in the absence of an input signal the cathode will not remain positive, but will return to the clamping potential. If the diode connections are reversed, the top of the wave will be clamped.

10-9 Cathode Follower.

Instead of the conventional amplifier, a "cathode follower" stage is sometimes used. The circuit is shown in Fig. 10-17. The plate load is the cathode resistor, and the output is taken across this. The output voltage subtracts from the input voltage, and the net input signal is small. In order to obtain any net input, the output voltage must be less than the input voltage, and so the gain of a cathode-follower stage is less than unity, or in other words it does not amplify.

Why, then, should a cathode-follower stage be used? The reason is that the output impedance is very low. This means that if a load is connected in parallel with the cathode resistor, the output voltage will hardly be changed at all. In other words the output voltage has very good regulation.

In an ordinary amplifier the equivalent circuit is as in Fig. 10-18 (a). The input voltage e_g is multiplied by the amplification factor μ , and applied to a series combination of the plate resistance R_p and the load resistance R_L . In the cathode follower R_L is replaced by R_K the cathode resistor. If i is the current flowing, a voltage drop iR_K is built up and this is applied in series opposition in the grid circuit. This may be taken into account as in Fig. 10-18 (b), by putting an additional resistance μR_K in the equivalent circuit, as the voltage drop is an input signal and must be multiplied by μ . This makes the total load resistance $(\mu + 1) R_K$, so we divide both voltage and resistance by $(\mu + 1)$ as in Fig. 10-18 (c). This will give the same current, by Ohm's Law, and so gives the correct output voltage. From this we see that the input voltage is multiplied by the factor $\frac{\mu}{\mu + 1}$, which is less than unity, and the internal resistance is divided by $(\mu + 1)$. If R_K is large in comparison with $\frac{R_p}{\mu + 1}$, the output voltage will be very nearly $\frac{\mu e_g}{\mu + 1}$, and will not be much affected by changes in R_K due to loading.

The cathode follower differs from an ordinary amplifier in another way. In the ordinary amplifier there is a phase reversal between input and output circuits, but in the cathode follower there is no such reversal.

Cathode followers can accept large input signals without running into grid current, because when the grid does go positive to earth the cathode also goes positive, and only the difference counts. With very large input signals clipping can however occur.

Since the gain of a cathode follower is less than unity, several stages in succession are not used. The cathode follower is only used when some special purpose is served, requiring good regulation of an output voltage.

10-10 Effects of Valve and Wiring Capacitance.

In sine-wave amplifiers stray capacitances in valves and wiring reduce the gain at high frequencies. With pulses the characteristic effect is the production of an opposite "spike" before a change followed by a sloping instead of a vertical line, and a rounded corner, as shown in Fig. 10-19 (a).

If the grid of a triode receives a positive pulse, a negative pulse occurs in the plate voltage. The pulse is applied across the grid-to-cathode capacitance, and in parallel with this is a capacitive voltage divider consisting of the grid-to-plate and plate-to-cathode capacitances in series, as in Fig. 10-19 (b). This will cause a positive spike in the plate voltage. Immediately afterwards the valve begins to pass current, and this begins to discharge the plate-to-cathode capacitance and causes the plate voltage to fall. Thus the spike is in opposition to the main pulse.

RADIO NAVIGATIONAL AIDS.

10-11 Voltage Transients in Parallel R, L, C. Circuits.

Inductances as well as resistors and capacitances are used in shaping circuits, particularly where narrow pulses of energy are required. Self inductances are frequently used in conjunction with capacitances in the generation of timing pulses, while mutual inductances in the form of pulse transformers are used at interstage couplings. Inductances and capacitances are used to build up artificial transmission lines used in controlling pulse durations and in producing pulses of special shape.

In R, C. circuits we are concerned with the time constant, but in R, L, C. circuits we are concerned with the resonant frequency of the combination, and on the Q of the circuit. Consider the circuit of Fig. 10-20 (a) and assume a steady current flowing in the inductance L. If now the switch is opened, and if the resistor R is open circuit, then the energy stored in L will cause a sine wave oscillation in L and C at their resonant frequency, as shown in Fig. 10-20 (b). An open circuit R will make the circuit equivalent to an L, C. circuit of infinite Q as it must be remembered that this resistor will in practice often be in part the internal resistance of the other circuit elements.

If now the same procedure is adopted but with R of a high value, then the opening of the switch will produce a damped oscillation in the circuit, as shown by the dotted line in Fig. 10-20 (b). The oscillation is now said to be damped due to the presence of R, which reduces the Q of the circuit. Energy dissipated in R represents a loss of energy in the circuit. The frequency of the oscillation will be slightly lower than in the case of the undamped circuit. If now the value of R is decreased, the rate of decay of the oscillation is increased. As long as the variations of voltage in the circuit remain oscillatory, the circuit is said to be underdamped. When however R is reduced to a value equal to $\frac{1}{2} L/C$, the damping is so rapid that only the first half-wave is produced as shown in Fig. 10-20 (c), and the circuit is said to be critically damped. Further decrease of R results in the overdamped wave shown in Fig. 10-20 (d).

In actual circuits, C may be a capacitor or just the self capacity of the inductance plus stray capacities in the circuit. A resistor may be placed in parallel with L and C if critical or overdamped response is required, or the damping may be only that caused by coil losses if a highly oscillatory transient is required.

10-12 Ringing Circuit.

One method of range measurement used in radar sets uses a resonant circuit in which an underdamped oscillation is begun at the time of each transmitted pulse. The oscillation continues for the duration of the range sweep and is damped out at the end of the sweep. In this way the range sweep is divided into a series of known intervals determined by the periods of the oscillation.

A typical circuit for producing such oscillations is shown in Fig. 10-21 (a). The name "ringing circuit" is applied to it because the oscillations it produces are similar to the mechanical oscillations of a bell when causing a ringing sound. To study the action of the circuit, it should be visualized in the normal condition,

with the valve conducting and a steady current flowing through L. At the instant of transmission the tube is cut off by a negative gate pulse, Fig. 10-21 (b), applied to the grid. The L.C. Circuit then "rings", and the oscillations, Fig. 10-21 (c) are produced. Losses in the L.C. circuit are kept as low as possible in order that the amplitude of the oscillations will be nearly constant throughout the duration of the gate pulse.

At the end of the gate pulse, the valve conducts once more, starting up a second oscillation. The impedance of the conducting valve however is in shunt across the L.C. circuit causing the oscillations to be rapidly damped out. It is usual after the production of the damped oscillations to shape it by clipping either the positive or negative half waves, and peaking the remainder, these two operations being carried out in subsequent parts of the circuitry.

10-13 The R, L, C. Peaker

The R, L, C. peaker circuit shown in Fig. 10-22 (a) is very similar to the ringing circuit. The capacitance of the circuit however is usually only the unavoidable self capacity of L and stray capacities. The resonant frequency is therefore high. A resistance R is connected across L, C. to produce nearly critical damping when the valve is cut off. Instead of an oscillation therefore, a single sharp positive peak is produced at t_1 Fig. 10-22 (b) and (c) corresponding to the beginning of the negative gate pulse, and a single sharp negative peak at t_2 , corresponding to the end of the pulse gate. The negative pulse is smaller than the positive pulse due to the extra damping afforded after t_2 by the conducting valve.

The inductance is connected in the anode circuit of the valve rather than the cathode circuit to obtain an output pulse of greater amplitude and with a steeper leading edge. The applications of this circuit are similar to those of the R.C. peaker, though use of the inductance facilitates the production of very narrow pulses of high amplitude.

10-14 Artificial Lines.

Inductances and capacitance may be used to build up an "artificial line", which has similar properties to a transmission line.

An ordinary transmission line has distributed inductance and capacitance, reckoned at so much per unit length. If ℓ is the inductance per cm length in henries and c of the capacitance per cm length in farads, the quantity

$R_0 = \sqrt{\frac{\ell}{c}}$ is called the "surge resistance", or characteristic resistance of the line, and is measured in ohms. The quantity $V = \sqrt{\frac{1}{\ell c}}$ is the velocity of propagation of a wave along the line, and is in cm. per sec. when ℓ and c are in the above units. The characteristic resistance is the input resistance of an infinite length of line. It may be used as a non-reflecting termination of a finite line, causing it to behave as if infinite.

A real line terminated in the characteristic resistance R_0 will transmit a

RADIO NAVIGATIONAL AIDS

Fig. 10-21

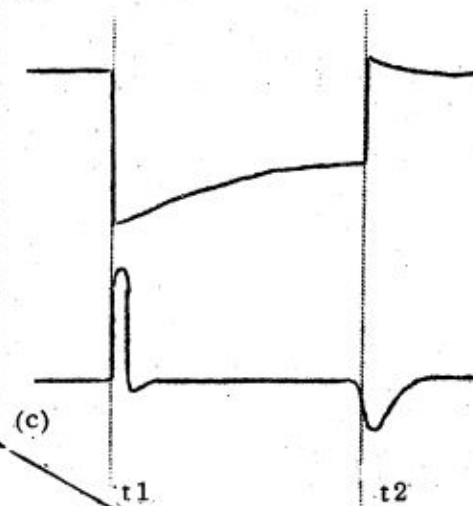
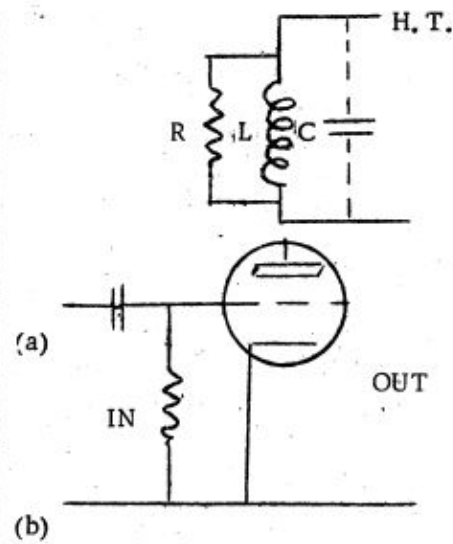
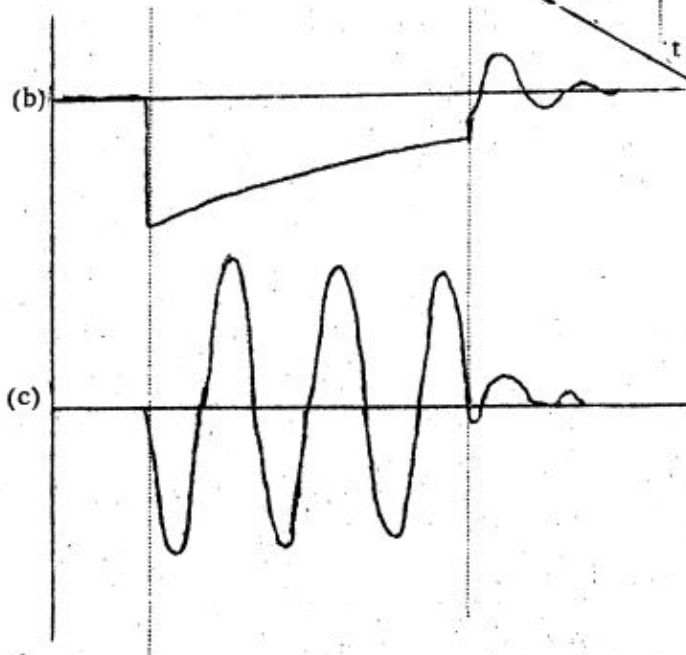
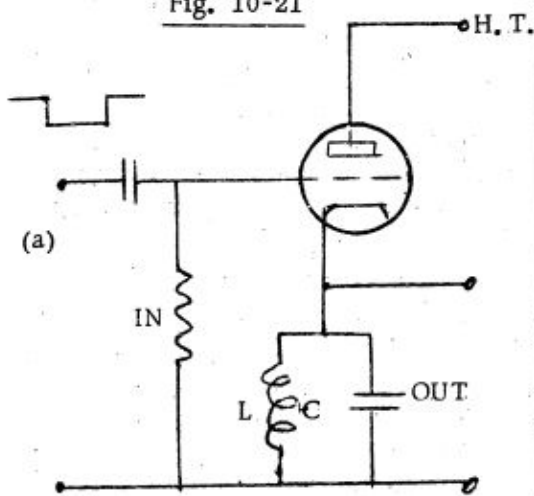


Fig. 10-22

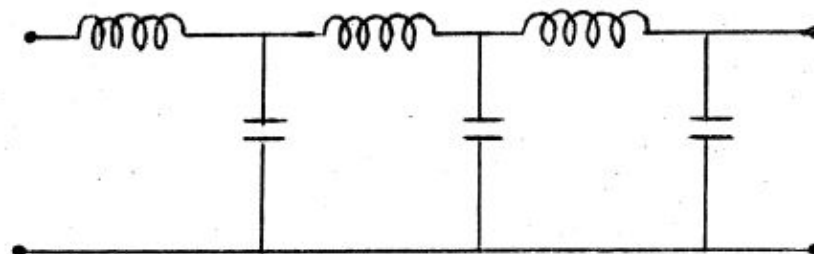


Fig. 10-23

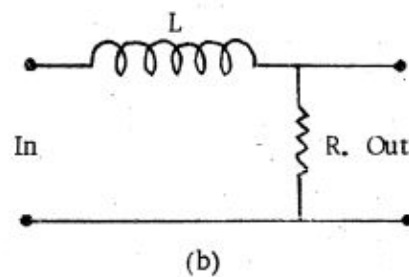
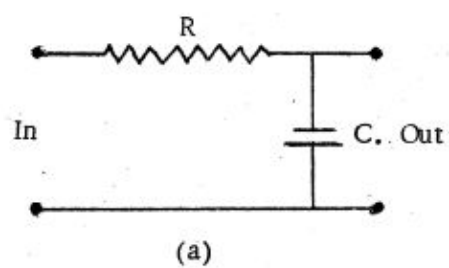


Fig. 10-24. Integrating Circuits

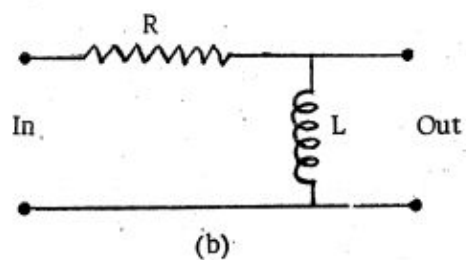
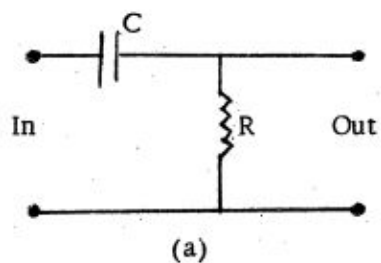


Fig. 10-25 Differentiating Circuits

RADIO NAVIGATIONAL AIDS.

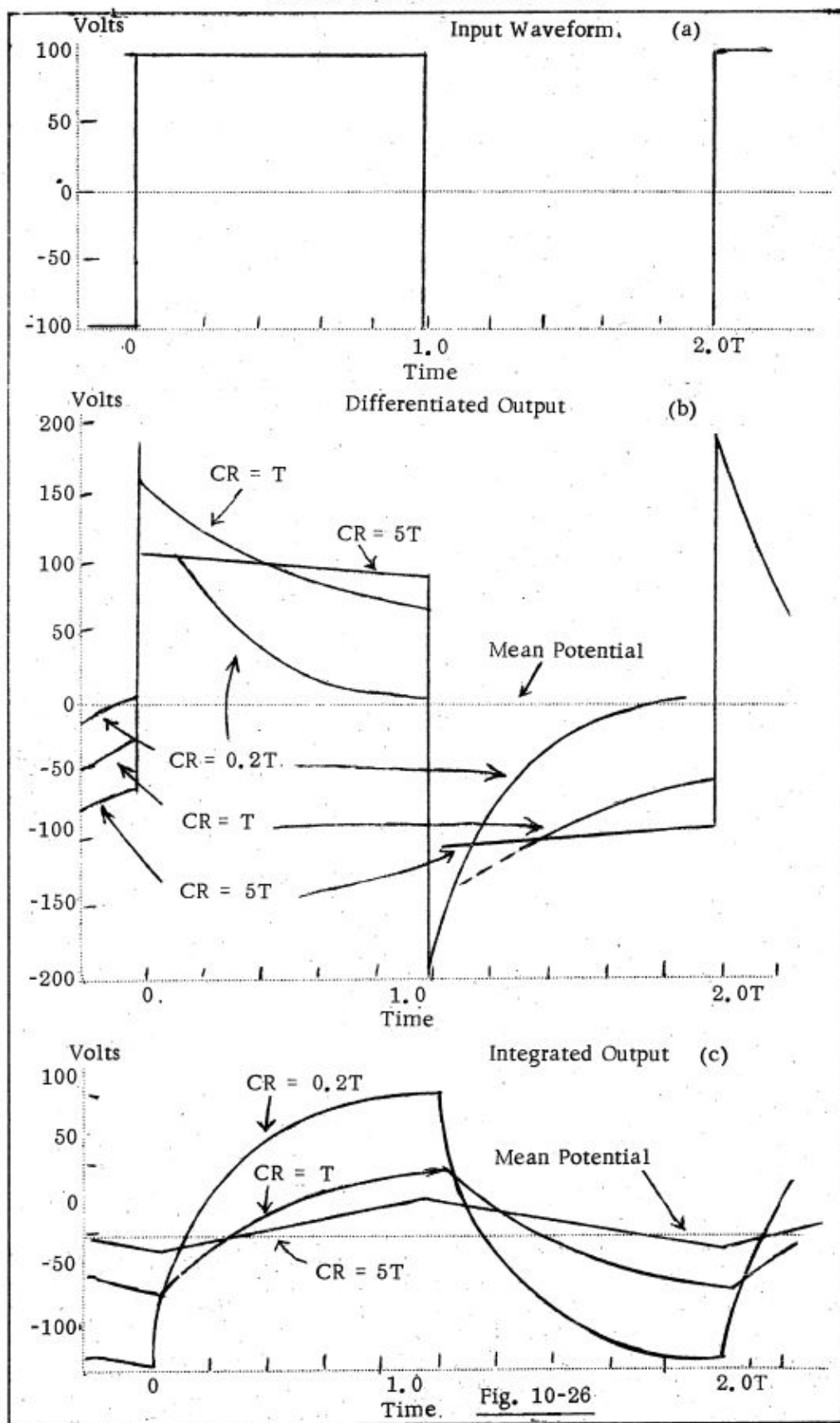


Fig. 10-26

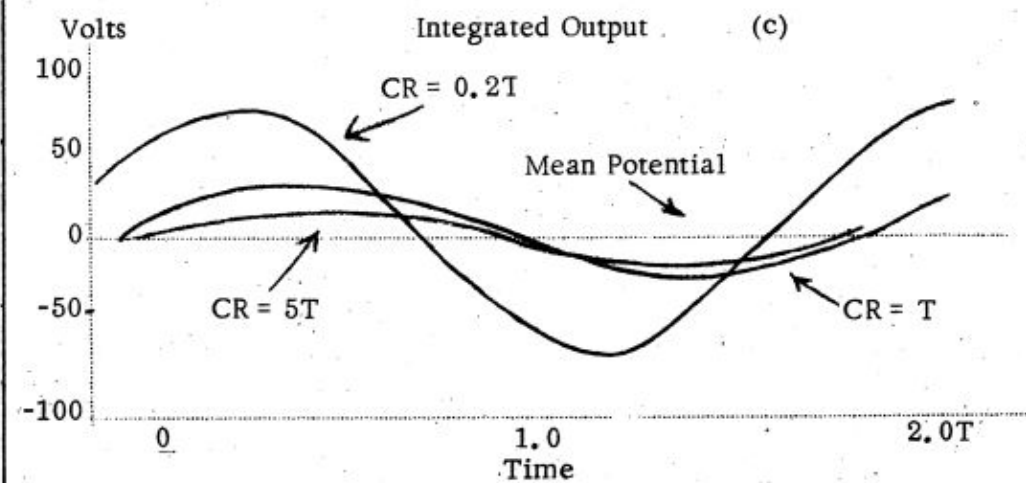
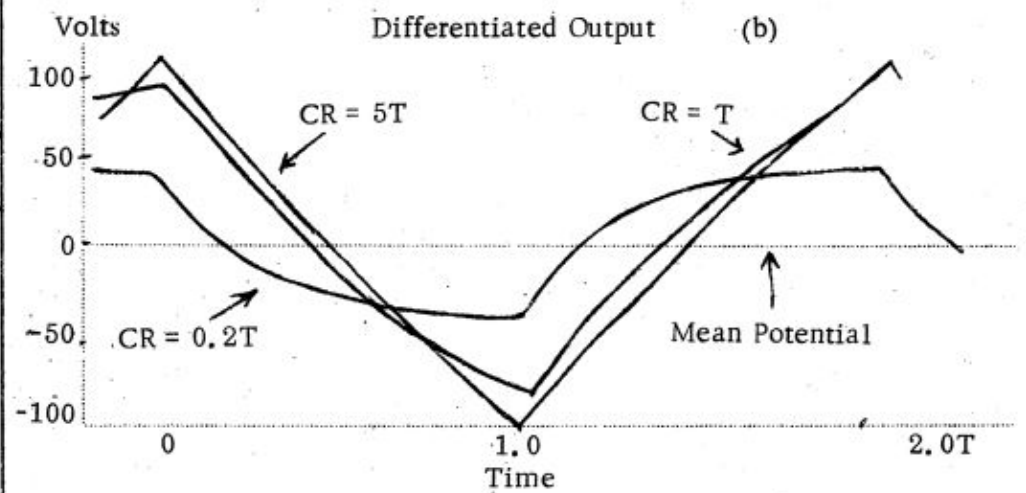
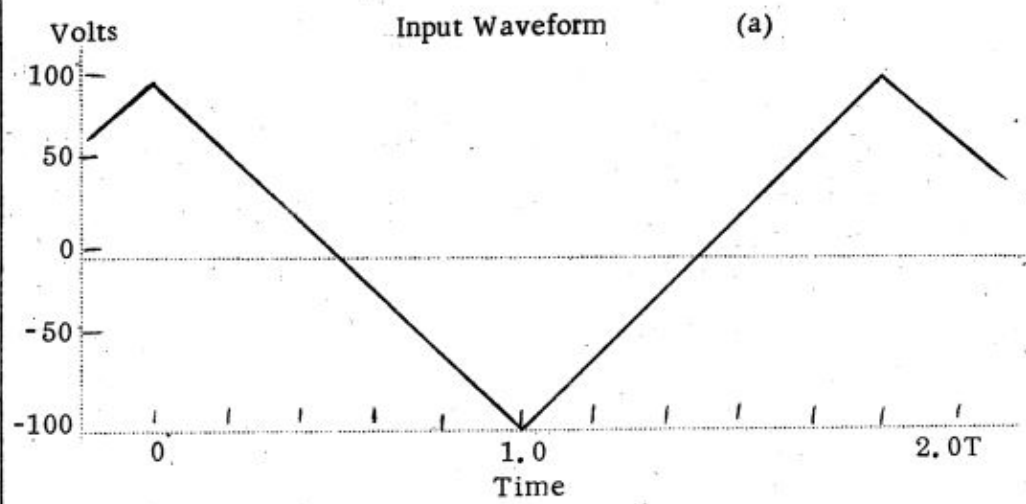


Fig. 10-27

pulse of any shape at the velocity v , without change of shape.

In practice real lines are used for transmission, but sometimes the effect of a line is required inside a piece of equipment. For example, a line 3000 metres long properly terminated will transmit a pulse from one end of the other in 1 microsecond, and it may therefore be used as a delay device. This is only a short delay, but the line is too long to build into equipment, so we use an artificial line.

An artificial line consists of a series of coils each representing the inductance of a suitable length of line, and a set of condensers each representing the capacitance of the same length of line, arranged as in Fig. 10-23. It behaves in practically the same way as a real line for low frequencies, where each section only represents a small fraction of a wavelength. If however the frequency is so high that each section represents more than about three-tenths of a wavelength, the artificial line ceases to behave as a line. Actually it is a low-pass filter with a definite cut-off frequency.

When a sudden pulse is applied to a real line, it behaves as a resistance of value R_0 . But an artificial line will at first behave as an inductance of capacitance, whichever is the first element. After a short interval it will behave as a resistance. Thus there will be a distortion of steep wave fronts. This distortion may consist of damped oscillations of a very high frequency.

If a transmission line, real or artificial, is charged to a suitable d-c voltage and is then suddenly discharged through a resistance equal to its characteristic resistance, the pulse of current will be square in form, and the pulse duration will be equal to the time taken for a wave to travel from one end of the line to the other and back. Further discussion of artificial lines and their applications will be found in Chapter 12 on Modulators.

10-15 Pulse Transformers.

Pulse transformers are used in the formation of pulses, particularly pulses of large amplitude, and also as inter-stage couplings in pulse generating circuits. For examples of the pulse transformer, reference should be made to Chapter 12, dealing with the subject of Modulators.

10-16 Summary Integrating and Differentiating Circuits.

Summarising the properties and the action of simple integrating and differentiating circuits, it will be useful to display first the circuits in question (see overleaf)

In Figs. 10-24 (a) to 10-27 (a), where the circuit elements are C and R , the time constant of the combination is CR seconds. In Figs. 10-24 (b) and 10-25 (b), where the circuit elements are L and R , the time constant of the combination is L/R seconds. For integration, the time constant of the circuit must be long compared with the duration of the applied pulse, whereas for differentiation, the time constant must be short compared with the duration of the applied pulse.

Fig. 10-26 demonstrates the integration and differentiation of a square waveform. Fig. 10-26 (a) shows the waveform, one half-cycle of which has a duration T seconds. Fig. 10-26 (b), shows integration using circuits of various time constants. In this case the degree of integration increases with increase in time constant. Fig. 10-27 shows integration and differentiation of a symmetrical sawtooth waveform. With these examples it should be possible for the student to deduce approximately the effect of these circuits on waveforms of various types.

- - - 0 - - -

QUESTION PAPER FOR CHAPTER 10.

1. Explain the clipping action of a triode.
2. What is the meaning and importance of the time-constant of a circuit ?
3. What is the difference in circuit values between a peaker and an amplifier.
4. What is the purpose of a clamper ? Describe one form of clamp.
5. Explain the action of the cathode follower.
6. Describe the effect of interelectrode capacitances upon the operation of a pulse amplifier.
7. Describe the action of a ringing circuit.

- - - 0 - - -

MARCONI SCHOOL OF WIRELESS
RADIO NAVIGATIONAL AIDS
CHAPTER 11 - PULSE TIMING.

11-1 Introduction.

In the previous chapter methods have been described for producing square waves and sawtooth waves, commencing with sine waves and using amplifying valves. The saw tooth waves require square waves to be produced first, and the duration of the rising part of the sawtooth is the same as the duration of the square pulse applied to the input. If it is desired to change the duration, corresponding to a change of range, it is necessary to change the length of the square pulse as well. It is also possible to generate square and sawtooth waves by means of oscillators. These oscillators are known as "relaxation oscillators", because the state of strain in a charged condenser is relaxed by discharging the condenser through a resistance. They fall into a number of distinct classes. Some will function without any input signal, but at the same time are capable of being "synchronised" with such a signal, so as to have both the frequency and the instant of commencing each cycle under control. Others respond to an input signal by performing one cycle of oscillation only and then stopping - these are expressively called "Flip-flops". Still others perform a half-cycle in response to a signal, and wait for another signal before completing the cycle. These are known as "Trigger Circuits".

Relaxation oscillators can also be divided into one-valve and two-valve types, and into hard-valve and soft-valve (gas-discharge) types. Soft valves include neon glow lamps and thyratrons. These are frequently used in time bases for cathode-ray oscillographs for testing purposes, but hard-valve types are used in radar. Soft valves have an internal "trigger" action, namely the "breakdown" or onset of ionisation at a definite applied voltage, and the extinction of the ionisation at a certain lower voltage. With hard valves, triggering is usually accomplished by cut-off, which may be produced by a negative potential applied to one of the electrodes. The difference between square-wave and sawtooth oscillators lies in the time-constants of the circuits. An abrupt change requires a short-time constant, a gradual change a long-time constant. In square-wave oscillators, a square wave of voltage may be obtained at one electrode and a sawtooth at another. As soft valve oscillators are not common in radar equipment, their treatment will be brief, and will be deferred until the end of the chapter.

11-2 Transitron Flip-Flop.

The behaviour of this circuit shown in Fig. 11-1 is best explained with the aid of a special set of characteristic curves, shown in Fig. 11-2. The control grid is set at a number of values of bias, one for each curve, and the suppressor potential is varied by connecting various positive or negative voltages in series with R_4 . For each setting there will be a certain screen current, and the drop across R_3 lowers the screen voltage. The suppressor voltage is plotted as abscissa, the screen voltage as ordinate.

For sudden pulses the condenser C_2 does not immediately alter its state of charge, and so the suppressor and screen voltages vary by the same amount, along a straight line such as ABC. With the circuit quiescent, the control grid and suppressor grid are both earthed, and so A is the operating point. If the circuit is disturbed by means of a small short pulse on the control grid, after the pulse the operating point must lie on ABC and on the curve for zero control grid voltage, and so at A, B, or C. If the operating point moves during the pulse to the right from A or C, it lies above the valve characteristic, and the higher screen voltage requires higher current, which can be drawn only from C_2 . This lowers the screen voltage, and the operating point returns to A or C respectively. A similar movement from B brings the operating point below the characteristic, and extra current is available to charge C_2 , since the screen current falls and the current in R_3 rises. The screen voltage must rise as C_2 is charged, and the operating point moves to A. In the same way, if the operating point moves to the left from A or C it will return, but if it moves to the left from B it will move to C. Thus B does not count as a stable state of balance.

Suppose that a positive pulse of say .15 volts is applied to the grid. The operating point is compelled to move along ABC, and must end on the curve for $E=+1.5v$ at the point C^1 . At the end of the pulse it must return to the zero-volt curve, but stops at C; with the suppressor biases strongly negative. This is the "flip". The plate current is now cut off, the plate voltage is high and the screen voltage low. Since there is no suppressor current, the suppressor cannot remain negative indefinitely, because of R_4 . Current flows in R_4 and partly neutralises the negative charge on the left plate of C_2 ; the corresponding positive charge on the right plate is neutralised by electrons collected by the screen. The screen at first remains at a fixed potential, until the point D is reached, at the sudden bend in the characteristic. The screen and suppressor now move together in a pulse (the "flop") to reach the point E. The suppressor is then positive and collects electrons, which accumulate on C_2 and reduce the voltage to zero, the conditions returning to the Point A. The cycle is now complete, and will not recommence until another pulse is applied.

Negative pulses also cause the circuit to operate. Suppose that a negative pulse of 1.5 volts is applied to the grid, then the operating point moves to the intersection A^- . This makes the suppressor positive. If the pulse lasts long enough, the condenser C_2 begins to charge and increases the difference between suppressor and screen voltages. This moves the operating line to the left of ABC,

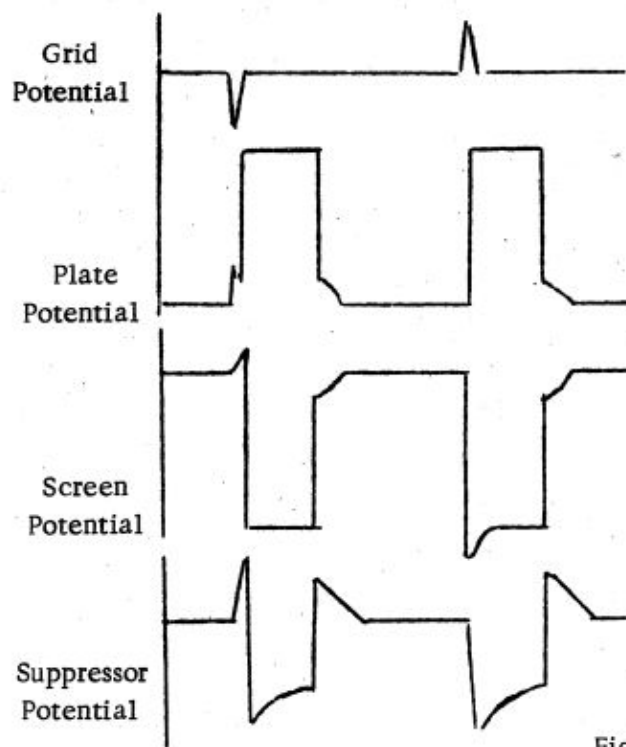


Fig. 11-3

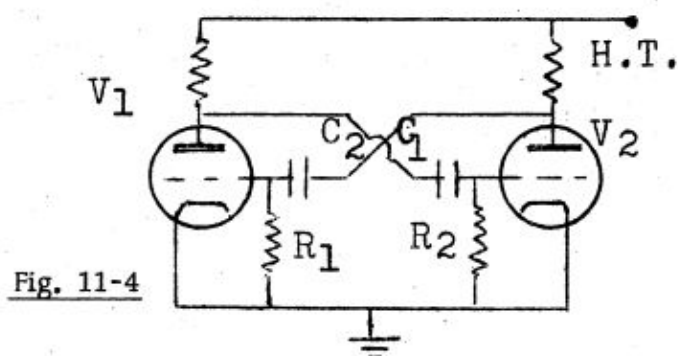


Fig. 11-4

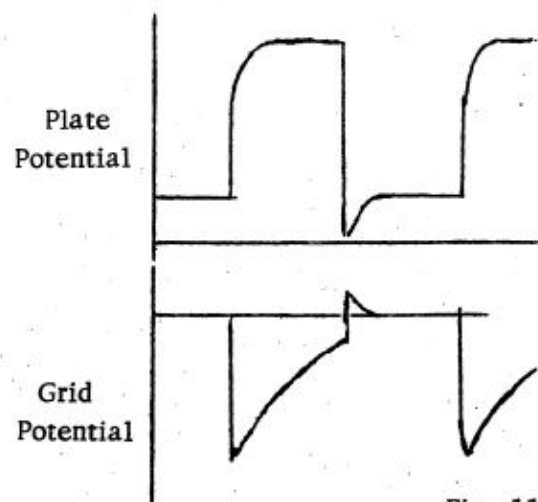


Fig. 11-5

RADIO NAVIGATIONAL AIDS.

and the new line may clear the bump between A and B. In this case, at the end of the pulse the operating point must move to a point on the zero control grid voltage curve, to the left of C. The cycle then continues as for positive pulses. Short pulses may fail to produce operation. The pulse should be kept small, because strong pulses may bring the point A into a region of reversed suppressor current, and cause the circuit to "Lock" with positive suppressor voltage, as C_2 discharges instead of charging. Operation then ceases.

The wave-forms for the two types of pulses are shown in Fig. 11-3. They differ slightly at the leading edge. Both plate and screen have nearly square waves. After the flop there is a short "tail" in all three waves.

11-3 MULTIVIBRATOR.

Oscillators using two valves are often arranged with each valve feeding back to the other. Such circuits are called multivibrators. They may be arranged with automatic triggering so as to produce oscillations continually, or they may be flip-flops.

The multivibrator circuit was first used as an oscillator for the production of waves rich in harmonics. It was synchronised to a source of known frequency, and used for frequency measurements. The original multivibrator circuit, due to Abraham and Bloch, is shown in Fig. 11-4. It consists of two triodes each with resistance-capacitance coupling to the other. The two valves should be matched, and it is convenient to use a twin triode. The circuit is usually made as symmetrical as possible.

Suppose that the two valves are carrying equal plate currents, then they will have equal plate voltages. If now the plate current in Valve V_1 increases for any reason, its plate voltage decreases, and a negative signal is applied, to the grid of the second valve V_2 , decreasing its plate current and raising its plate voltage. This passes a positive signal to the grid of V_1 , raising the plate current further. The current in V_1 rapidly rises to a maximum, and the current in V_2 is cut off. The second grid is then strongly negative, because of the charge on C_1 . This charge gradually leaks away through R_2 , until the grid voltage of the second valve reaches the cut-off point. This causes plate current to commence in the second valve, and a similar change-over occurs, so that the first valve is cut-off. The charge on C_2 then leaks away through R_1 until plate current commences.

The wave forms for the voltages on the grids and plates are shown in Fig. 11-5.

The frequency of repetition depends on the time constants $C_1 R_2$ and $C_2 R_1$ and on the valve characteristics. If the frequency is to be variable, C_1 and C_2 can be controlled in steps by ganged switches, and R_1 and R_2 may be ganged variable resistances.

If $C_1 R_2$ and $C_2 R_1$ are made unequal, or if one valve is biased, the two halves of the cycle will be of unequal length. The multivibrator may be synchronised by applying pulses to the grid of either valve, or to the two grids in opposite senses. Such pulses may be applied through a condenser, and are con-

veniently obtained from a "peaker" valve.

11-4 Multivibrator Flip-flops.

A flip-flop can be obtained by biasing one of the valves beyond cut-off, as in Fig. 11-6. This is most readily done by connecting the cathode of the valve concerned to a tapping on a bleeder resistance across the H. T. supply. The valve V_1 is then normally "on", and V_2 normally cut off. A small accidental decrease of the plate current in V_1 transfers a positive pulse to the grid of V_2 , as previously, but this does not cause any plate current in V_2 , and so there is no cumulative action and no change-over.

The operation may be started by a negative pulse on the grid of V_1 , or alternatively by a positive pulse on the grid of V_2 , sufficient to cause plate current. The negative pulse on V_1 reduces its plate current, raises the plate voltage and transfers a positive pulse to the grid of V_2 . Owing to the gain of V_1 as an amplifier, this will be larger than the initial pulse on V_1 . If the resulting pulse on V_2 is large enough to start plate current, feedback will occur as before, and V_2 will build up a large current, while V_1 will become cut off. The condenser C_2 then discharges until the plate current of V_1 starts again. The circuit then flops back to the original condition with V_2 cut off, and remains in this stage, until another pulse occurs.

The wave-forms of the various electrode voltages are given in Fig. 11-7. Square waves of both polarities are available, at the grid and plate of the cutoff valve, while a sawtooth wave is available at the grid of V_1 , the valve normally on. The duration of the sawtooth wave can be controlled by C_2 and R_1 .

A modification of the above circuit uses pentodes, the pulses being applied to the suppressor of V_1 in a negative direction or of V_2 in a positive direction.

The behaviour with regard to pulses of either sign is different from that in the transitron. In the multivibrator opposite pulses must be applied to different electrodes, while in the transitron flip-flop they may be applied to the same electrode.

Fig. 11-8 shows another multivibrator flip-flop circuit, known as "cathode-coupled" type; Fig. 11-6 is said to be "plate-coupled". The circuit resembles the ordinary multivibrator circuit, but the coupling condenser C_2 is omitted, and instead, both cathode currents flow through a common resistor R_5 . The resistance R_5 should be about double R_4 . The grid leak R_2 , which is of very high resistance is returned to H. T. Valve V_2 therefore normally has a slight positive bias, depending on the grid current, and passes plate current, which produces a voltage drop across R_5 , sufficient to bias V_1 to cut-off.

The circuit is started by a positive pulse on the grid of V_1 , large enough to cause current to flow. This produces a negative voltage pulse on the plate of V_1 , which is transferred to the grid of V_2 and reduces the plate current of this valve. This lowers the voltage drop across R_5 and so lessens the bias of V_1 . The current in V_2 is rapidly cut off. The condenser C_1 then discharges through R_2 and R_3 , until V_2 begins to pass current. This causes a voltage drop in R_5 , which acts as a negative pulse on the grid of V_1 and reduces its plate current.

RADIO NAVIGATIONAL AIDS.

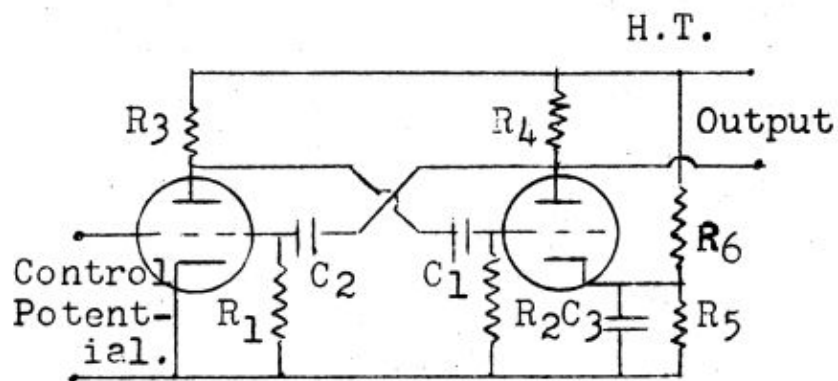


Fig. 11-6

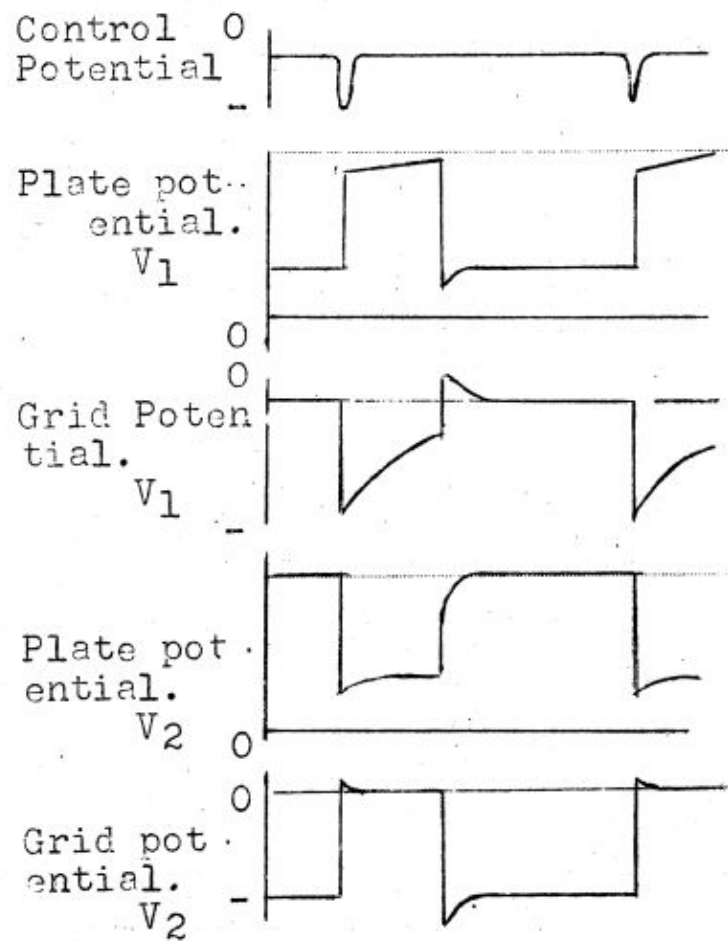


Fig. 11-7

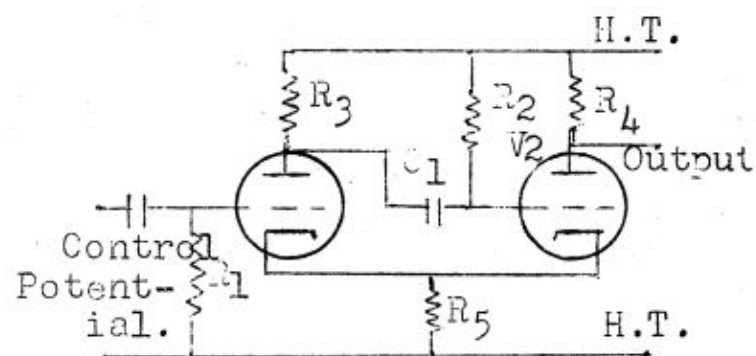


Fig. 11-8

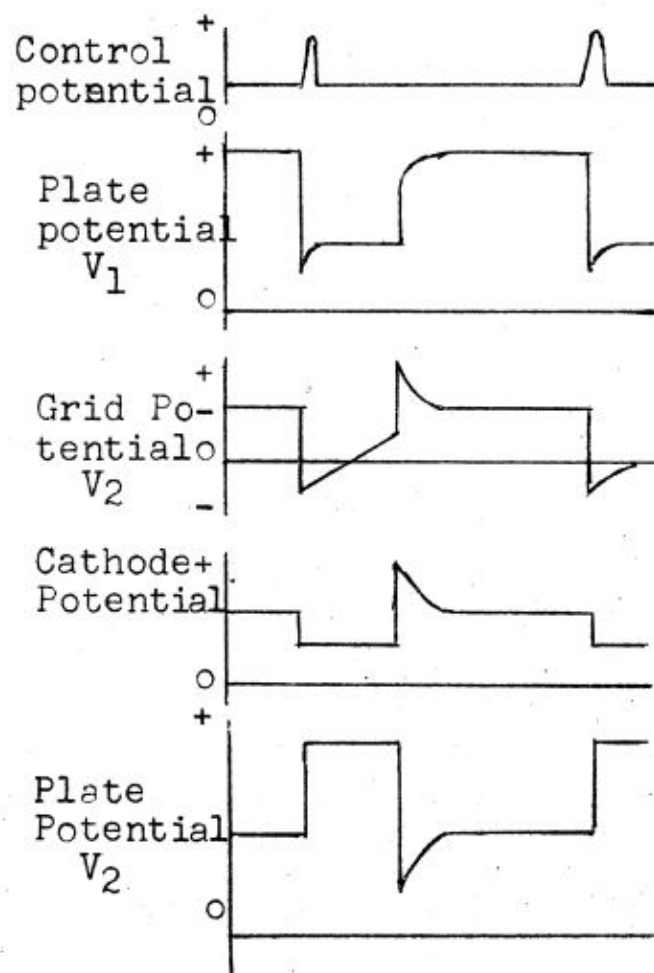


Fig. 11-9

RADIO NAVIGATIONAL AIDS

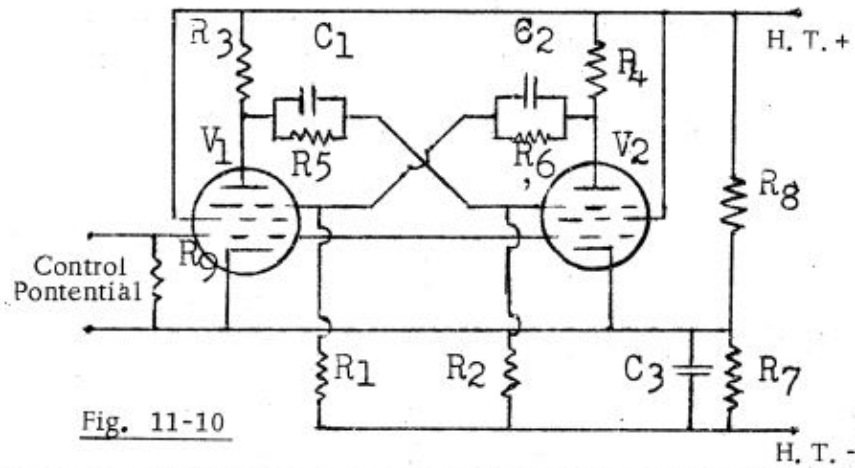


Fig. 11-10

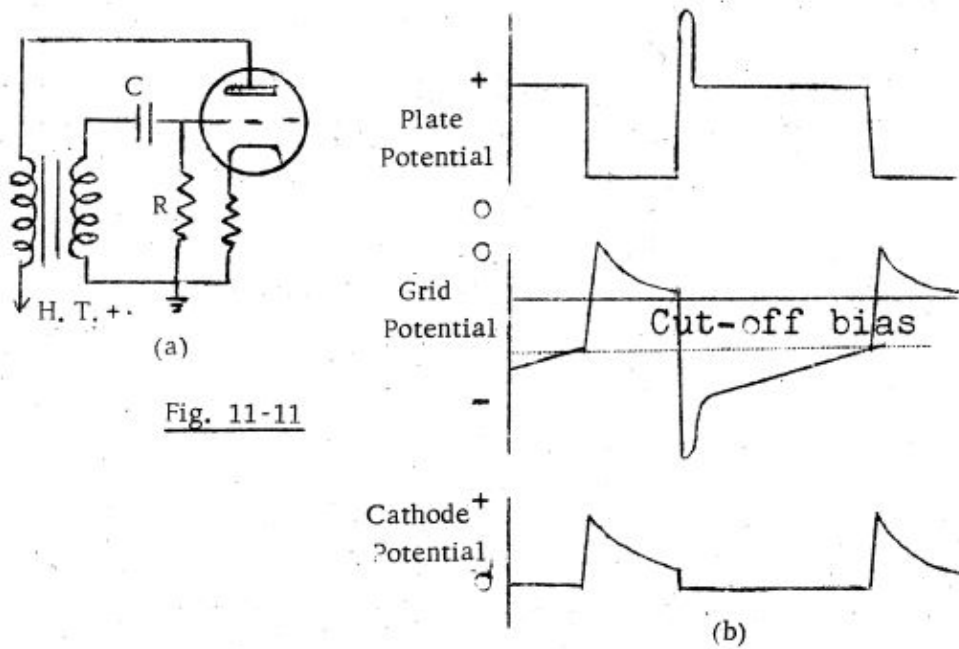


Fig. 11-11

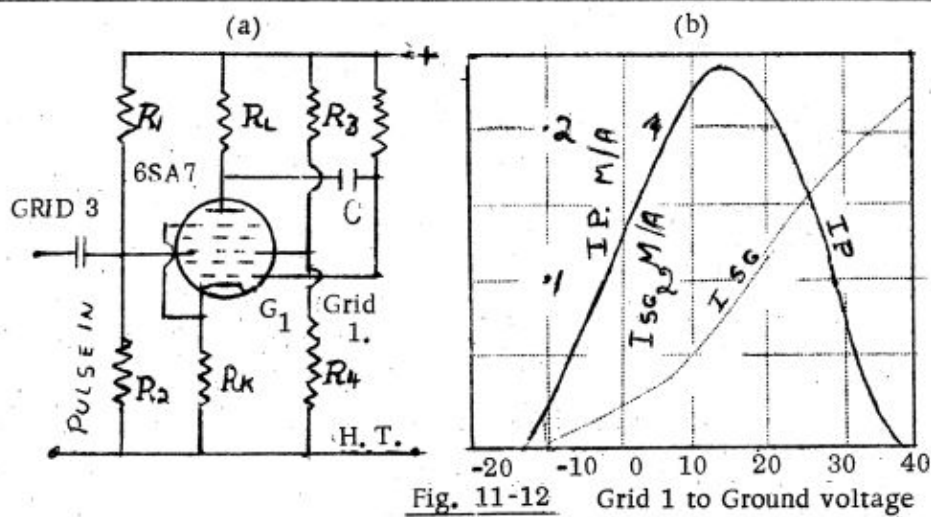


Fig. 11-12

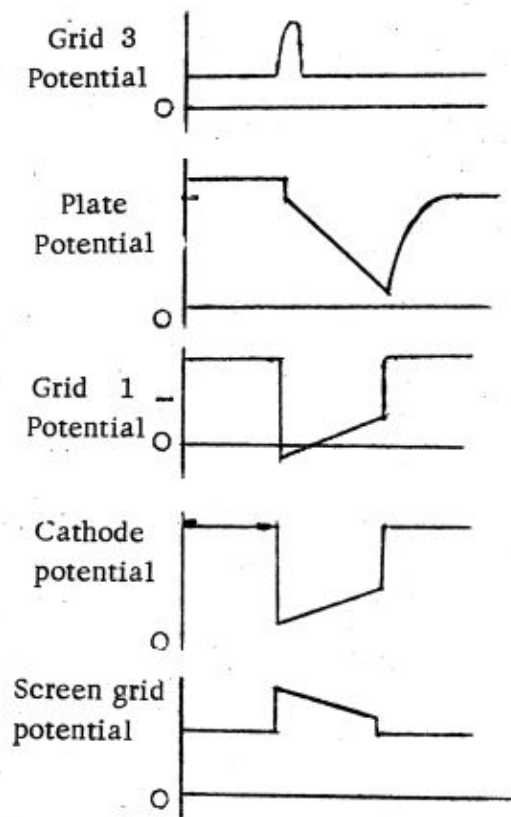


Fig. 11-13

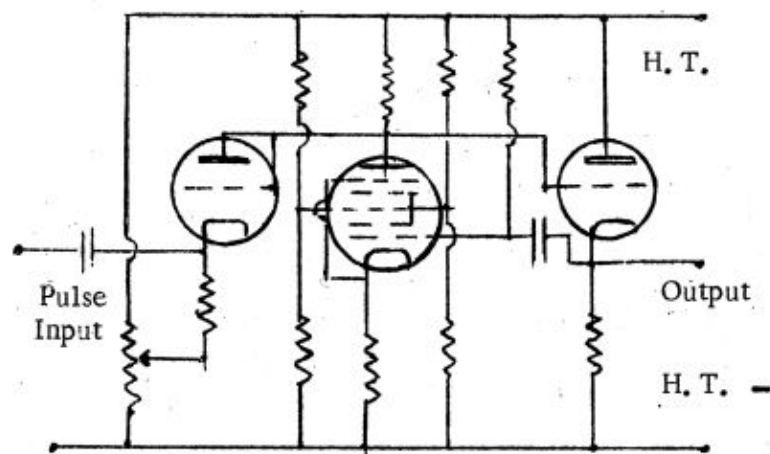


Fig. 11-14

RADIO NAVIGATIONAL AIDS.

The rise of plate voltage in V_1 is transferred to the grid of V_2 and increases the current. V_1 is thus cut off again and the cycle is complete.

The wave-forms of the various voltages to earth are shown in Fig 11-9

11-5 Eccles-Jordan Circuit.

A circuit originally due to Eccles and Jordan, but somewhat modified later, is shown in Fig. 11-10. This is similar in appearance to the multivibrator except that the coupling condensers are shunted by resistors. The performance is however quite different.

The multivibrator proper has two unstable states, with either valve cut off. The flip-flop can remain in one state until it is triggered to the other, which is unstable, so that it returns, and the oscillation stops after one cycle. The present circuit can remain with either valve cut off, but can be triggered from one state to the other. The oscillation therefore stops after one half-cycle, and a second pulse is required to complete the cycle.

In the improved circuit shown, couplings are taken to opposite suppressors instead of to control grids as in most multivibrators. Negative pulses are applied to both grids in parallel to trigger the circuit. Bias is applied to the suppressors only, by means of a bleeder across the H. T. supply, to which the cathodes are attached at a suitable point.

The intervalve couplings by the resistors amount to voltage dividers determining the suppressors voltages as fractions of the opposite plate voltages. The condensers are added to pass steep wave fronts. The time constants of the coupling circuits should be approximately equal to the duration of the pulses applied to the grids.

Suppose that V_1 is cut off and V_2 is conducting. A negative pulse on the grids reduces the plate current of V_2 and raises its plate potential. This causes a positive pulse on the suppressor of V_1 , sufficient to cause current flow to begin. This lowers the plate potential of V_1 and consequently also lowers the suppressor voltage of V_2 . The valve V_2 then triggers off, and this should happen at exactly the peak of the applied grid pulse. As the grid returns to normal the plate current in V_1 rises to its full value. A second pulse will cause V_1 to cut off and V_2 to take current.

11-6 Blocking Oscillator.

Very short and regular pulses may be obtained by what is known as a "blocking oscillator". In an ordinary oscillator provided with a grid leak and condenser, the bias developed builds up rapidly to an equilibrium value. In the blocking oscillator the bias builds up to cut-off, in the first cycle. The oscillator starts, generates a pulse, and stops till the charge leaks off the grid condenser sufficiently to cause a restart.

A typical circuit is shown in Fig. 11-11 (a). The feedback is through an iron-cored "pulse" transformer, the secondary of which is coupled to the grid through a condenser, the grid leak being earthed at the bottom. A small resistance may be placed in the cathode lead to develop an output voltage, or

alternatively output may be taken from the plate by resistance-capacitance coupling to a second valve.

When the H. T. is switched on, the grid is at zero potential and the valve will draw current. This lowers the plate voltage and a positive voltage is induced in the secondary and applied to the grid. The transformer usually has equal numbers of Cathode turns in the two windings, and so the rise potential in grid voltage is equal to the fall in plate voltage. The result is a further rise in plate current, so that the plate voltage rapidly becomes low and the grid voltage high, as in the upper diagram of Fig. 11-11 (b).

Grid current flows as well as plate current when the grid becomes positive and this lowers the grid voltage by charging the condenser. The plate current is therefore reduced and the plate voltage rises. A negative voltage is therefore applied to the grid through the transformer and condenser, and the grid is rapidly driven beyond cut-off. A strong positive pulse then appears in the plate voltage, and a negative pulse in the grid voltage. The plate voltage returns to normal when the plate current ceases changing (being zero) and the grid voltage gradually rises by discharge of the condenser, until the cut-off voltage is reached. The cycle then recommences.

The voltages from plate, grid and cathode to earth are shown in Fig. 11-11 (b). Square negative pulses and short positive pulses are obtained at the plate, while a sawtooth voltage is available at the grid. The duration of the sawtooth depends on the grid leak and condenser while the duration of the period of conduction depends mainly on the transformer characteristics, especially the leakage inductance.

The circuit can be converted into a flip-flop by returning the grid to a point of negative potential, so biasing it just beyond cut-off. A positive pulse on the grid will then produce one cycle of operation.

The free-running oscillator without control pulses is useful as a master oscillator for pulse work, the short positive pulses on the plate being used to trigger other circuits. These pulses may be made less than 0.3 micro-seconds in duration without very great difficulty.

11-7 Phantatron Circuit.

A circuit that may be used for controllable delay is known as the phantatron. The basic circuit is shown in Fig. 11-12 (a). A pentagrid converter valve, such as the 6SA7, is used in a circuit very similar to a cathode-coupled multivibrator. No. 3 control grid and the plate act as valve V_1 , while No. 1 control grid and the screen (grids 2 and 4) take the place of V_2 . No. 5 grid is a suppressor.

There is however a difference from the two-valve multivibrator. The plate current corresponding to V_1 can be cut off by control No. 3, but the screen current cannot be cut off without stopping the plate current. Thus the changeover does not quite cut off the screen current.

The action of the circuit may be explained with the aid of the characteristic curves of Fig. 11-12 (b) in which plate and screen currents are plotted against

RADIO NAVIGATIONAL AIDS

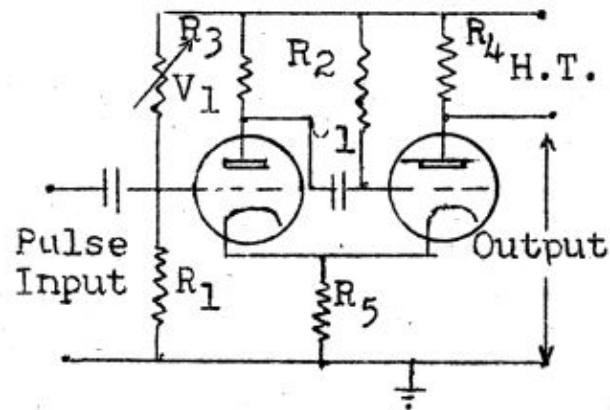


Fig. 11-15

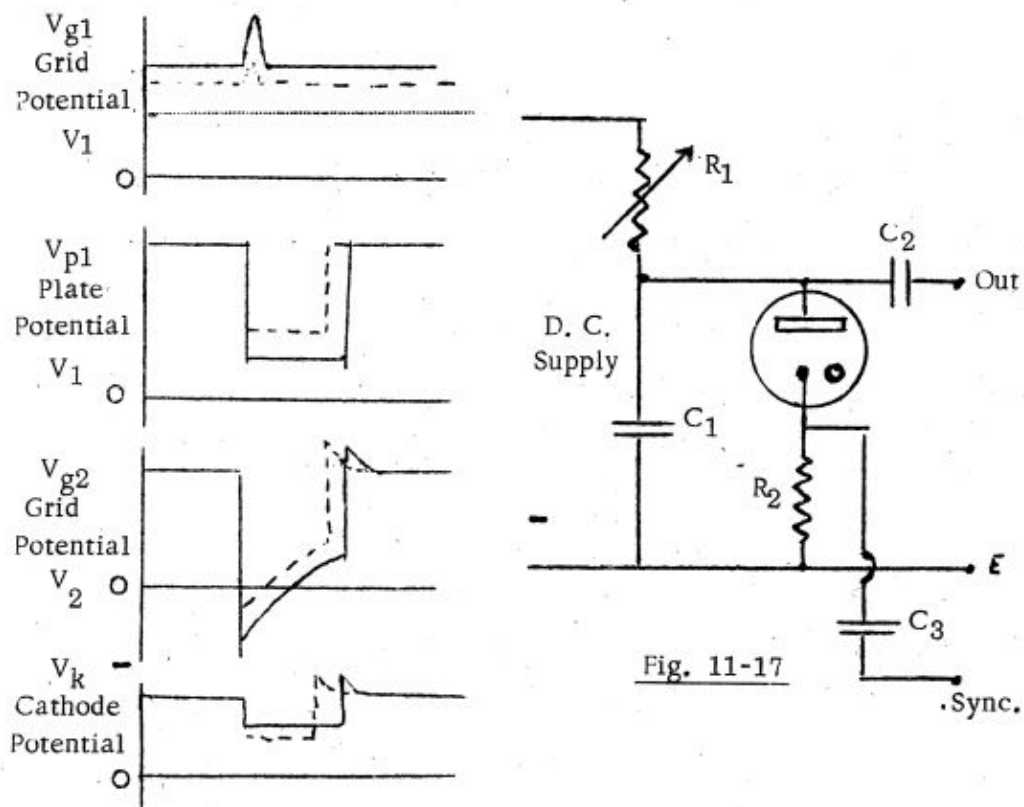


Fig. 11-17

Fig. 11-16

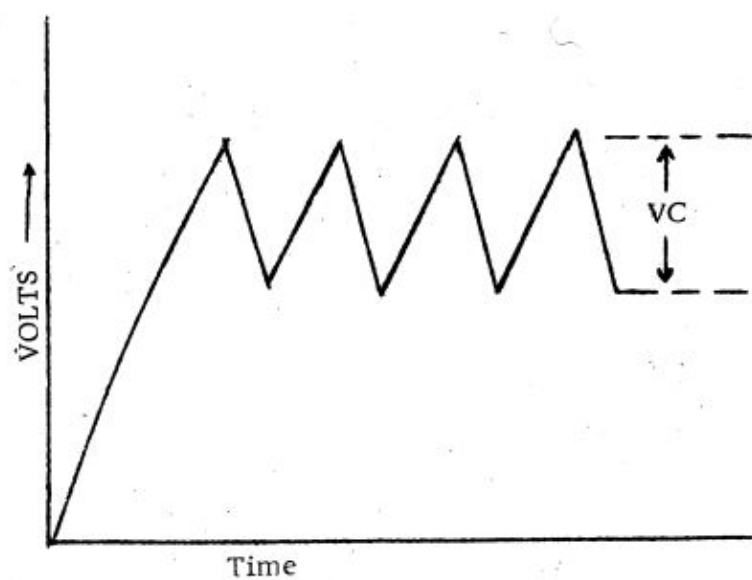


Fig. 11-18

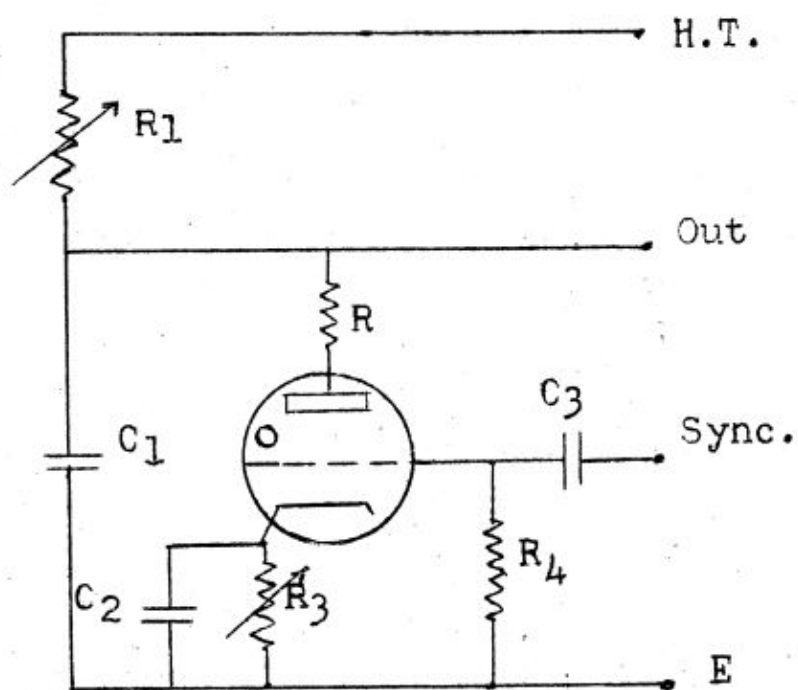


Fig. 11-19

RADIO NAVIGATIONAL AIDS.

the voltage of No. 1 grid to earth, with No. 3 grid at a fixed voltage to earth. When No. 1 grid is strongly negative all current is cut off. As it becomes less negative, the screen and plate currents rise, and the voltage from cathode to earth also rise. This rise tends to bias No. 3 grid negatively, so that the plate current reaches a maximum and falls, while the screen current still increases. At a certain positive voltage on No. 1 grid, relative to earth, the plate current is cut off. It is usually arranged that the No. 1 grid is then practically at cathode potential, so that grid current is about to commence. If No. 1 grid is returned to H. T. through a high resistance grid leak, it is held at cathode potential between pulses.

Before the arrival of the input pulse, the No. 1 grid is at cathode potential and the screen current is high. When a positive pulse is applied to grid No. 3, plate current begins to flow and the plate voltage falls. Grid No. 1 is coupled to the plate by the condenser C, and its potential falls. This causes the plate current to rise and the screen current to fall. The plate current rises to a maximum and falls to a position of balance. The screen current is initially larger than the plate current, and so the total cathode current has fallen, so that the cathode potential also falls.

After the input pulse is over the condenser C begins to discharge, allowing the No. 1 grid potential to rise as in Fig. 11-13. This causes an increase of plate current, and the plate potential falls much faster than the No. 1 grid voltage rises. The cathode potential gradually rises. After the peak of the plate current curve is reached, the plate current begins to fall, and the plate voltage begins to rise. This raises the No. 1 grid voltage, and a second switching action occurs. The voltages of the cathode, No. 1 grid, and the screen return rapidly to normal, but the plate voltage recovers only slowly as the condenser recharges through the plate-circuit load resistance.

The voltage curves of Fig. 11-13 show a down-going sawtooth for the plate voltage, of considerable slope, and also up-going sawtooth waves of very small slope for No. 1 grid and cathode.

The performance may be improved by adding a cathode follower to aid in recharging the condenser, and a diode to aid in controlling the pulse length, as in Fig. 11-14, in which the diode and cathode follower are the two halves of a twin-grid triode such as the 6SN7, with grid and plate joined together for the diode. The diode cathode is tapped on to a potentiometer across the H. T. supply, and the setting controls the length of the phantastron pulse.

A negative trigger pulse is applied to the diode cathode, causing it to conduct, and so lowering the pentagrid voltage. This change is passed on to the triode, grid and its cathode potential falls, so that a negative pulse is applied to the No. 1 grid of the phantastron, thus starting the change-over. The condenser is recharged from the cathode follower, so that the long curve in the plate voltage is removed. The diode bias limits the charging of the condenser and so controls the pulse length. The trailing edge of the pulse is not very sharp, and the output, taken across the cathode follower load resistance, is usually applied to trigger a blocking oscillator and gives a sharp delayed pulse. The delay may

be read from the dial of the diode-biasing potentiometer.

11-8 Multivibrator Delay Circuit.

The cathode-coupled multivibrator circuit given in Fig. 11-8 may be modified to give controllable delay. The delay is controlled by biasing the grid of the normally off valve V_1 . This is conveniently done by connecting the grid of V_1 to a voltage divider across the H. T. supply as in Fig. 11-15. The tap must be so arranged that the voltage from grid to earth, although positive is less than that from cathode to earth, by an amount great enough to cause cut-off of V_1 . If V_1 is not cut off, the circuit will not operate as a flip-flop. On the other hand if the bias is too large the circuit cannot be triggered.

In the normal state V_1 is cut off, V_2 conducts. One end of condenser C_1 is at H. T. potential, the other practically at cathode potential. After the input pulse, V_1 is conducting and V_2 cut off. The plate current of V_1 depends on the voltage applied to its grid, and so the fall of plate potential will also depend on the grid voltage. The cathode potential depends on the current in R_5 , and falls because V_1 passes less current than V_2 did. Fig. 11-16 shows the wave forms of voltages to ground, for two different values of applied grid voltage.

For the lower value of grid voltage the plate current of V_1 , is smaller, and thus the plate voltage of V_1 does not fall so far as for a higher grid voltage. On the other hand, the cathode voltage falls lower. The initial fall in grid voltage of V_2 is equal to the fall of plate voltage of V_1 , and is hence smaller for the small grid voltage. The grid condenser then discharges to the cut-off point of V_2 . For the smaller grid voltage of V_2 , the cathode potential is lower and the grid of V_2 does not need to rise so high to reach the cut-off point. The initial charge of C_1 is the same in both cases, and the cut-off point will therefore be reached earlier, with the smaller value of V_{g1} . Thus the duration of the conducting period of V_1 is under control. The positive jumps in V_{g2} and V_k depend on V_2 , and are similar in the two cases.

With proper relative proportions of R_3 and R_5 , if the duration of conduction of V_1 is plotted against V_{g1} a straight line is obtained. The positive jump in V_k is then delayed by an amount that can be readily found by measuring V_{g1} , or by reading the dial controlling V_{g1} .

11-9 The Neon Relaxation Oscillator.

A neon tube is a two electrode valve, filled with neon gas at low pressure and in which neither electrode is heated. When a potential, known as the striking potential, is applied across the electrodes, the gas becomes ionized, and the impedance of the valve falls from a very high to a very low value. If now the voltage applied is lowered, the extinction potential of the valve is reached, about 70% of the striking potential, and the valve deionizes. The impedance rises once again to a very high value.

The neon saw-tooth oscillator of Fig. 11-17 consists of a condenser C_1 charged through a high resistance R_1 , until the potential on C_1 becomes the striking potential of the neon tube. At this point, the neon strikes and C_1 is

RADIO NAVIGATIONAL AIDS

discharged through R_2 and the low impedance of the ionized tube, until C_1 potential falls below the extinction potential of the tube. The device then repeats this operation indefinitely. Fig. 11-18 shows the voltage variations at the output, and it will readily be seen that the amplitude of the output wave is limited to the difference between the striking and extinction potentials of the tube in use. The repetition frequency can be controlled by variation of R_1 or C_1 , while synchronisation can be achieved by the application of negative going pulses to the "sync" terminal.

11-10 The Thyatron Time Base.

The principle of operation of the Thyatron time base is very similar to that of the neon oscillator. The thyatron valve has large separation between striking and extinction voltages, and furthermore the separation can be easily controlled. The thyatron is a triode valve, with heated cathode, and containing a small quantity of inert gas. The striking voltage depends mainly on the grid potential, and when held off by grid bias, the valve behaves like a hard valve biased to cut-off. On the approach of the anode voltage to the striking voltage, a small anode current passes, partly ionising the gas. A cloud of positive ions surround the grid, causing it to lose control. Once this has happened, the impedance of the valve falls to a very low value, the voltage drop across the valve being about 15, largely irrespective to the current flowing. When the gas is ionised, it can normally be deionized only by the reduction of anode voltage, and not by reduction of grid voltage. Grid potential can thus be used to control the commencement of ionization, but not to control cessation.

A simple type of sawtooth generator is shown in Fig. 11-19 and its resemblance to the neon oscillator will be recognised immediately. As in the case of Fig. 11-17, the repetition frequency can be controlled by variation of R_1 or C_1 . An additional variable control, R_3 , is provided, which functions as an output amplitude control. The time constant of $R_3 C_3$ should be large compared with the time duration of the longest pulse the circuit is required to produce. In this way an increment of charge is passed into C_2 each time the valve fires, building up a final bias voltage on the valve. The repetition frequency and amplitude controls in this type of circuit are to some extent interdependent, movement of one usually requiring readjustment of the other. A resistor R_2 must always be used in the anode circuit to limit current in the valve to a safe figure. Synchronisation can be achieved by the application of positive going pulses to the grid of the valve.

QUESTION PAPER FOR CHAPTER 11.

1. With the aid of suitable diagrams describe the action of the Transitron flip-flop.
2. Describe the action of the multivibrator. Describe the modification necessary to convert the circuit into a flip-flop circuit ?
3. Describe the action of the blocking oscillator.
4. Describe a circuit capable of introducing a controllable measure of delay.
5. Describe the operation of a thyatron time base.

MARCONI SCHOOL OF WIRELESS

RADIO NAVIGATIONAL AIDS.

CHAPTER 12 - MODULATORS

12-1 Introduction.

In plate modulation of telephony transmitters, the plate voltage is varied in accordance with the speech wave, and the power input consequently varies. For 100 percent modulation the peak power rises to four times the average power.

Plate modulation is also used in radar. The H. T. voltage is applied for a very short period, and then removed for a much longer interval. The power during the pulse is many times greater than the average power. Suppose for example that we apply 10,000 volts to a load of 1000 ohms, then the power is

$$\frac{E^2}{R} = \frac{10^8}{10^3} = 10^5 \text{ watts or } 100 \text{ kw.}$$
 If the pulse lasts for 1 micro-second, the energy in the pulse is $W = Pt = 10^5 \times 10^{-6} = .1 \text{ joule}$. If there are 1000 pulses per second, the average power is $.1 \times 1000 = 100 \text{ watts}$. Thus the average power is only one-thousandth of the peak power. This ratio is known as the "duty cycle" and is equal to the pulse length in seconds multiplied by the repetition frequency.

Different radar transmissions may differ in carrier frequency or in pulse repetition frequency. The cathode-ray tube circuit is synchronised with the pulse frequency, and the receiver is tuned to the carrier frequency. Suppose that two transmitters are operated on the same carrier frequency, but have different pulse frequencies. The received signals will be amplified and detected, and applied to the cathode ray tube. The wanted signals produce a stationary pattern, but the unwanted signals produce successive traces in different parts of the screen, and their net effect is to produce a diffused illumination, which alters the contrast between the desired signal and the background. The interfering signal does not prevent reception. This corresponds to the reception of a CW signal in wireless telegraphy in the presence of others by listening only to sounds of a certain pitch.

The power developed by the transmitting valve is greatest when the applied voltage is greatest, and falls rapidly as this is reduced. Thus for maximum output the voltage should be maintained at its maximum value for as great a portion

MARCONI SCHOOL

of the pulse as possible. The wave form of the pulse should have a flat top and steep sides; vertical sides cannot be produced.

A radar modulator then consists of an arrangement to produce pulses of short duration, and to transform or amplify them to give the required voltage. Arrangements must be made to repeat these pulses at regular intervals.

12-2 Simple Pulse-forming Network

Pulses may be produced by charging a condenser through a resistance from a d-c supply, and discharging it through the load by suddenly closing a switch. This switch may take the form of a spark gap, which becomes conducting when the applied voltage reaches a certain value, and ceases to conduct when the current falls to a low value insufficient to maintain ionisation.

The circuit of Fig. 12-1(a) is the most obvious arrangement. The charging resistance R is much larger than the load resistance R_L . The condenser C charges with time constant $C(R + R_L)$, and the charging current begins at

$\frac{E}{R + R_L}$ and falls to a low value. The load voltage has the initial value $\frac{ER_L}{R + R_L}$, which is only a small fraction of E . When

the breakdown voltage is reached, the gap is short-circuited, and the condenser discharges. If E_m is the maximum condenser voltage, the voltage across R_L will jump to the same value in the opposite direction, since the sum of the two voltages equals the gap voltage, which is taken as zero. The voltages across C and R_L decrease together, with time constant CR_L , and the current in this branch of the circuit falls correspondingly. The gap also carries during discharge a current equal to $\frac{E}{R}$ from the source, and when the total current is small the gap will

open. The charging resistance R must be large enough to prevent a permanent discharge through the gap. When the gap opens the cycle recommences.

Two voltage wave-forms are of interest, the condenser voltage, which is shown in Fig. 12-1 (b) and the voltage across the load, shown in Fig. 12-1 (c). The condenser voltage changes gradually but the load voltage has sudden jumps at the beginning and end of the discharge. The values between pulses are small being at most equal to $\frac{ER_L}{R + R_L}$, while the negative peak value at the commencement of discharge is nearly equal to the breakdown voltage of the gap.

The pulses may be applied to the cathode of a transmitting valve, the plate of which is earthed. A negative pulse then corresponds to a positive plate voltage.

The above system suffers from three defects. The pulse shape is poor, the pulses are irregular, depending on the properties of the gap, and a great part of the input power is wasted in the charging resistance. Pulse shape may be improved by using an artificial transmission line instead of a condenser, and regular pulses may be obtained by special spark gaps that fire in response to a control pulse. The power loss may be reduced by charging through an inductance.

RADIO NAVIGATIONAL AIDS

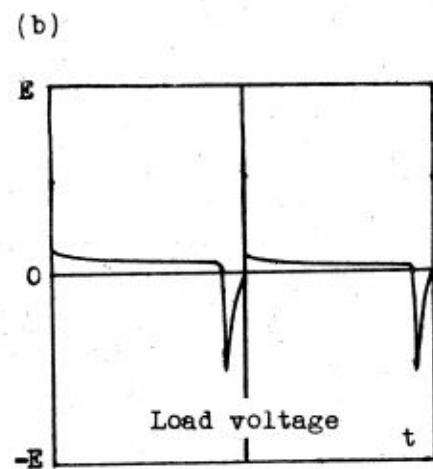
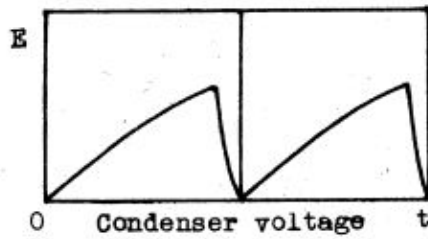
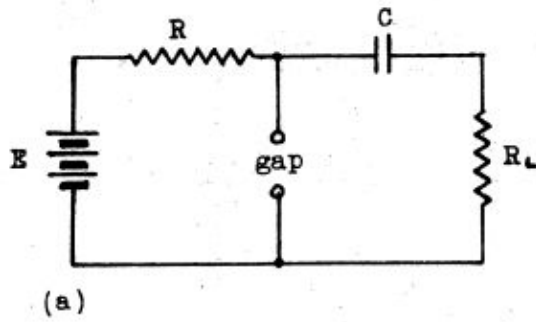


Fig. 12-1

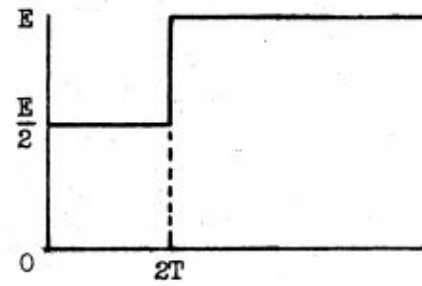
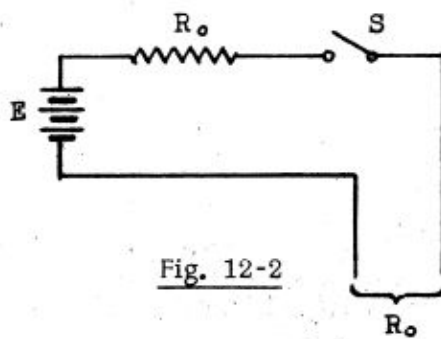


Fig. 12-3

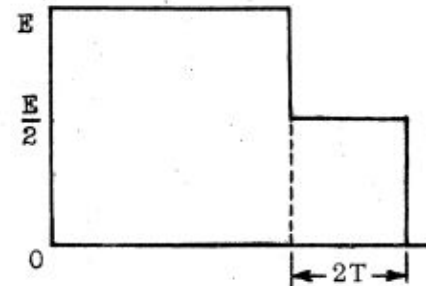


Fig. 12-4

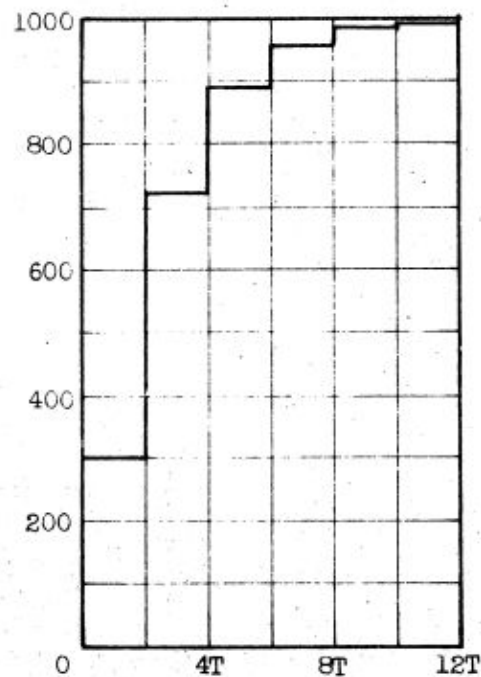


Fig. 12-5

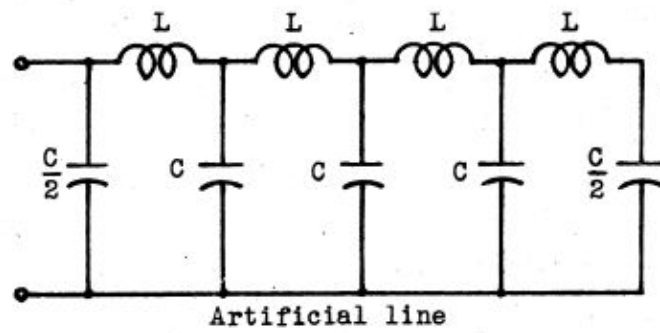


Fig. 12-6

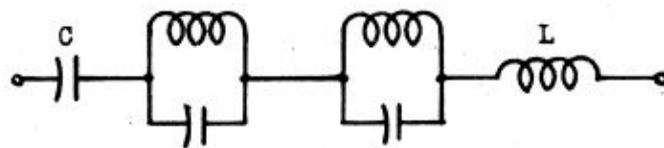


Fig. 12-7

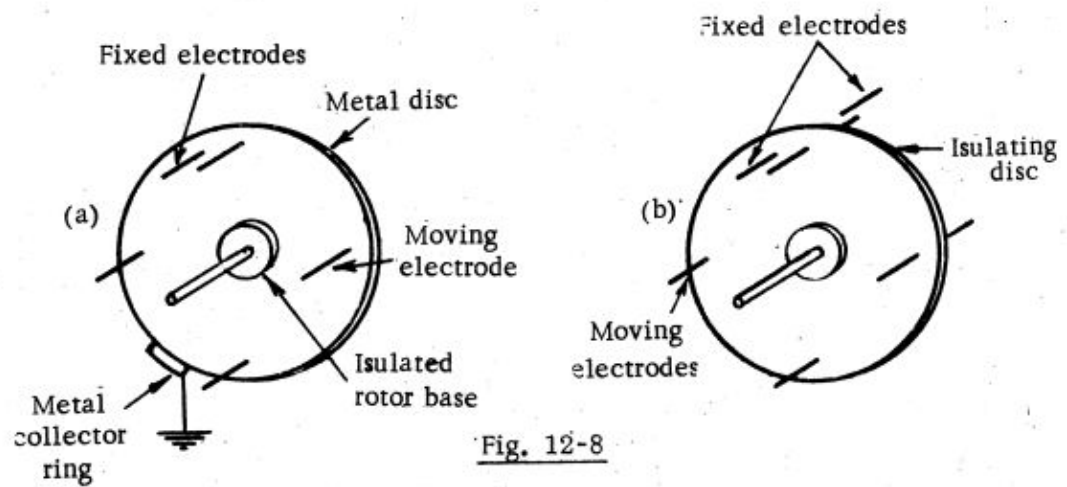


Fig. 12-8

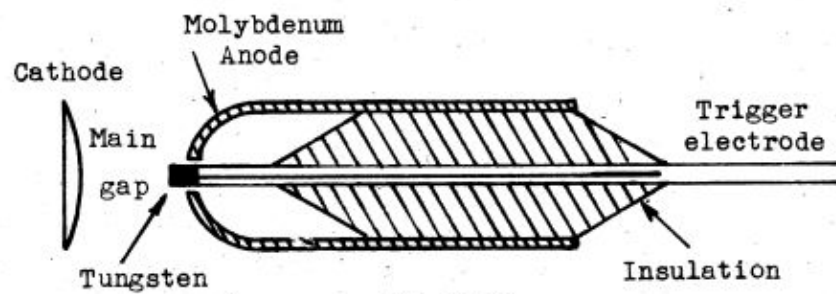


Fig. 12-9

RADIO NAVIGATIONAL AIDS.

12-3 Pulse Formation with Open Lines.

Pulses may be produced by charging an open-circuited transmission line instead of a condenser. The total resistance in the charging circuit should be made equal to the surge resistance R of the line. The charging circuit is shown in Fig. 12-2. The line behaves as a resistance R at the moment of closing the switch S , just as an infinite line. The voltage across the near end of the line is then $\frac{E}{2}$, and as current flows into the line a pulse of voltage moves along at uniform speed, until the end is reached. The whole line is then charged to voltage $\frac{E}{2}$. A reflection occurs at the open end, and in the return wave the voltage is equal and of the same sign, while the current is reversed. The line thus becomes charged to voltage E from the far end, while the current is zero from the open end as far as the return wave has reached. When the wave has returned to the input end, the current has ceased everywhere. The voltage at the near end is built up in two equal steps, as shown in Fig. 12-3. The time taken for complete charging is equal to twice the time for a pulse to travel from one end of the line to the other. The switch S may be opened after charging is complete. The line is allowed to remain fully charged for some time, and is then discharged by suddenly connecting a resistance equal to R_0 across the input end. The voltage immediately drops to half its initial value, and maintains this value for an interval equal to the time taken for a wave to travel to the end of the line and back, when it suddenly drops to zero, as in Fig. 12-4. In practice the load will have some inductance, which will prevent the instantaneous jumps, so that the sides of the pulse are steep, but not vertical.

If the line is charged or discharged through a resistance other than R_0 , the process will be carried out in a large number of steps of diminishing magnitude. Suppose for example that we charge a 600-ohm line through a resistance of 1400 ohms the applied voltage being 1000 volts. At the moment of closing the switch, the line behaves as a pure resistance, and the total resistance of the circuit is 2000 ohms. The current is .5 amp, and so the voltage across the input end of the line is 300, or 30 percent of the supply voltage. A wave travels to the far end of the line and is reflected totally, so that when it reaches the near end the voltage is 600. The difference between this and the supply is 400 volts and the circuit behaves as if 400 volts were suddenly applied. A new wave of 120 volts, or 30 percent of 400 volts, is produced at the input end of the line, raising the total to 720 volts. When the reflected wave returns, the line is charged to 840 volts. This is 160 volts below the applied voltage. A new wave of 30 percent of this, or 48 volts, is then sent down the line and reflected, raising the total voltage to 936 volts. Each new input wave is 40 percent of the previous one. The voltage of the input end of the line thus increases in steps, the successive values being 300, 720, 888, 955, 982 etc. Each of the values lasts for twice the time of transmission of a wave along the line. These values are plotted in Fig. 12-5. The discharge occurs in similar steps if the load resistance is different from R_0 . The charge and discharge are both complete in a single interval if the resistance is matched to the line.

If the line is charged or discharged through a resistor smaller than R_0 each step is too large, overshooting the mark, and successive steps are in opposite directions and of decreasing magnitude. When discharging, the line voltage becomes alternatively positive and negative in successive steps.

12-4 Pulse Forming Networks.

To obtain a pulse of 1 microsecond, the total distance of travel to the far end and back must be 300 metres, or the line must be 150 metres long. Such a line cannot be built into equipment, and so a compact circuit is used to imitate a line. One suitable circuit is an "artificial line" composed of a number of pi sections, as in Fig. 12-6. The series arms contain equal condensers C , except for the first and last, which have only half the capacitance, $\frac{C}{2}$. The surge resistance is given by $R_0 = \frac{L}{C}$, and the equivalent length of each section is found from the fact that the inductance per metre length of a parallel-wire line is $\frac{R_0}{300} \mu\text{H/m}$. The length per section is thus $\frac{300L}{R_0}$ metres, where L is in microhenries. The required length of line determines the total inductance and the total capacitance, but the number of sections can be varied. The greater the number of sections, the closer is the behaviour to that of a real line.

Instead of an artificial line, a network such as Fig. 12-7 may be used. It consists of a series combination of a condenser C , an inductance L , and several parallel LC circuits, tuned to a series of harmonic frequencies. Such a network has a reactance varying with frequency in nearly the same manner as the input reactance of an open-circuited line. At frequencies such that the line is a multiple of a half wavelength long, the input reactance is infinite if losses are neglected. These frequencies are made equal to the resonant frequencies of the successive parallel L-C circuits. By a suitable choice of component values, the reactances at other frequencies can be made nearly correct. The calculations are too long to be given. The frequencies chosen should represent fairly well those appearing in a pulse of the required duration, thus for a 1 μ sec pulse they should extend up to at least 1 mcps.

12-5 Discharge Devices.

Instead of allowing the charging circuit to cause the spark when a particular voltage is reached, more regular operation can be obtained by making the spark occur at regular intervals. Several methods are available for this purpose.

12-6 Rotary Gap.

The simplest method is to use a rotary gap, such as was formerly used for spark telegraphy. In some circuits one side of the gap is earthed. In such a case the rotary gap can consist of an earthed rotating disc carrying a number of contact points, which pass very close to a fixed contact point, as in Fig. 12-8 (a). The breakdown voltage of the gap is relatively small, so that the spark is certain to occur then, while the sparking distance of the main disc is too large for sparking at any other time. For best results the earth connection should be made by means

RADIO NAVIGATIONAL AIDS.

of a brush making contact with the disc, instead of using the shaft and bearings of the driving motor.

In circuits in which the gap is insulated the simplest arrangement is an insulated disc carrying pins which project on both sides and pass close to two fixed points, one on each side, as shown in Fig. 12-8 (b). This arrangement is flexible because one of the fixed points may be earthed if desired, so that such a gap can be used in any circuit. A spark in air produces traces of oxides of nitrogen and of ozone, which have a corrosive action. It is desirable to remove these gases from the neighbourhood of the gap by means of a pump.

12-7 Triggered Gap.

More satisfactory action can be obtained by means of a gap that is caused to act by means of an external impulse. A triggered gap that operates in the open is shown in Fig. 12-9. It consists of a molybdenum cathode, and a molybdenum anode which is in the form of a cup with a small hole at one end. A trigger electrode passes down the centre and terminates in the small hole in the anode. The details of supports are not shown.

The circuit builds up a large voltage between cathode and anode, but the actual spark is caused by a triggering pulse applied between the trigger electrode and the anode. This produces a small spark, and sufficient ions to reduce the breakdown voltage between anode and cathode, so that the main gap fires.

The triggering pulse may be obtained from the circuit of Fig. 12-10. This consists of a triode biased to cut-off and having its plate connected to an H. T. supply through a choke and also connected to the trigger electrode. A square pulse is applied to the grid in the positive direction, and builds up a current in the inductance. When the pulse is cut off, the grid becomes sufficiently negative to stop plate current, and this induces a very high voltage across the inductance. This voltage is applied to the trigger electrode, and breaks down the gap.

12-8 Trigatron.

The corrrison that occurs in ordinary open gaps can be avoided by using a special enclosed gap, in a glass vessel, known as a "trigatron". It is filled with a low pressure of inert gas, and contains two electrodes. Two gaps are generally used in series. The triggering voltage, derived in the same way as before, is applied between earth and the common point between the two trigatrons, as in Fig. 12-11, and breaks down one of them. This throws the whole of the charging circuit voltage across the other gap and breaks it down. It is found desirable to connect a high resistance across each gap, so as to equalise the drops across them when in the non-conducting state. These resistances are above one megohm each, and have no serious effect on the charging process. Trigatrons usually operate at about 5000 volts, and a pair can handle about 125 kw. Both rotary gaps and trigatrons are somewhat irregular in operation, and the intervals between pulses may vary by several microseconds. This variation is called "jitter".

12-9 Hydrogen Thyratrons.

The most satisfactory discharge device is a hydrogen-filled triode valve

known as a thyatron. In a triode there is a definite cut-off grid bias voltage for each value of plate voltage. If the grid bias is applied first, and the plate voltage built up afterwards by charging a line or condenser, no current will flow in the valve as long as the grid bias exceeds cut-off.

If now a positive pulse is applied to the grid so as to cause conduction, the hydrogen will become ionised, and the positive ions will be attracted to the negative grid. The grid is thus surrounded by positive charge which neutralises its effect. The gas takes up a potential practically equal to that of the plate, and the voltage drop to grid potential is concentrated in a thin sheath around the grid. A similar sheath forms around the cathode. The main part of the gas is free from electric field, and contains equal number of positive ions and electrons in any volume, so that it is free from space charge. There is an intense positive space charge in the sheaths over grid and cathode.

The result is that the grid can start the conduction as a result of the positive pulse, but then the gas is completely shielded from the grid, which loses control of the current. Flow of current will continue until the plate voltage is reduced below that necessary to keep the gas ionised. The sheaths then disappear and the grid resumes control. In the conducting state the current is limited only by the circuit resistance.

Mercury-vapour thyratrons have been in use for a much longer time than the hydrogen type. The advantage of hydrogen is that this is a permanent gas, and the warming-up period required with mercury is eliminated. The valve is ready for operation as soon as the filament becomes hot enough for emission. Thyratrons can replace other discharge devices in all the circuits given above, if means are provided to give the necessary positive grid impulse.

12-10 Charging Systems.

Several methods are available for charging the pulse-forming network or artificial line.

12-11 D-c Resistance Charging.

The circuit of Fig. 12-1 is an example of d-c resistance charging. It is not used for high power, because of the low efficiency. If the capacitance of the condenser is C farads, then CE coulombs are required to charge it to voltage E . This charge is obtained from the supply, and the energy drawn from the supply is therefore CE^2 joules. But the energy stored in the condenser is only $\frac{CE^2}{2}$ joules, and the remainder is dissipated in the charging resistance. Thus the efficiency of charging is only 50 percent.

If the condenser is replaced by a line or pulse-forming network, the time of charging through a resistance is long in comparison with the time of transmission through the network, and so the latter behaves as a condenser.

12-12 D-c Resonance Charging.

Fig. 12-12 (a) shows a more efficient arrangement for charging. An inductance L is used instead of a resistance, and the application of the voltage sets up

RADIO NAVIGATIONAL AIDS

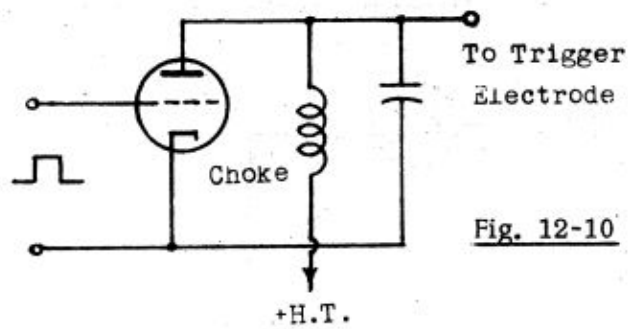


Fig. 12-10

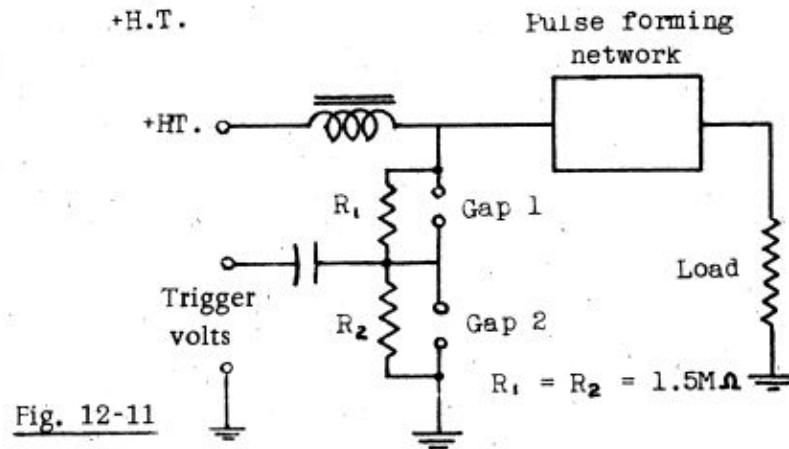


Fig. 12-11

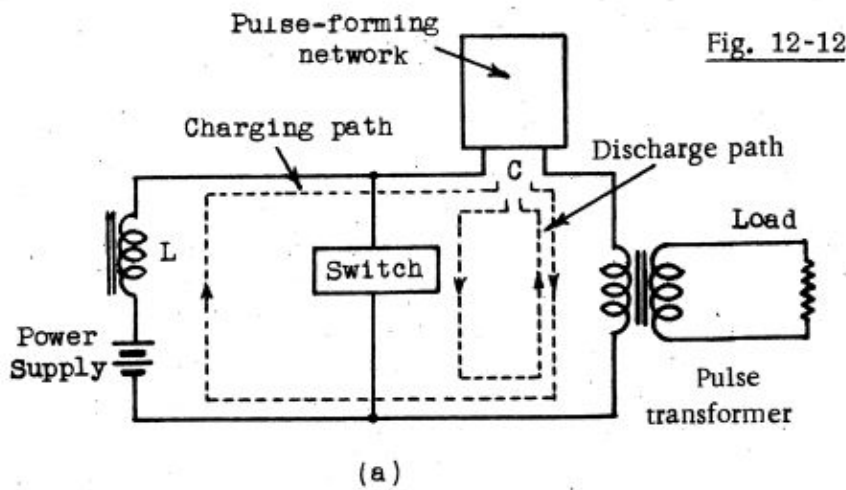
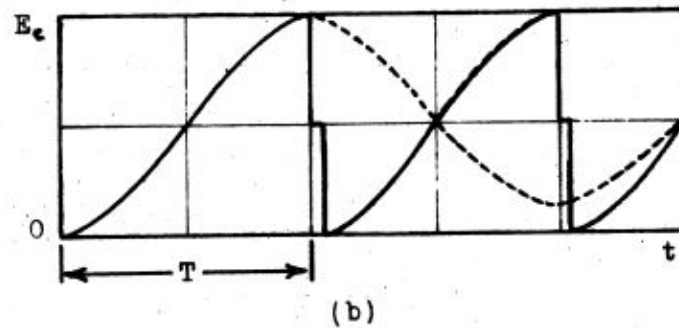


Fig. 12-12



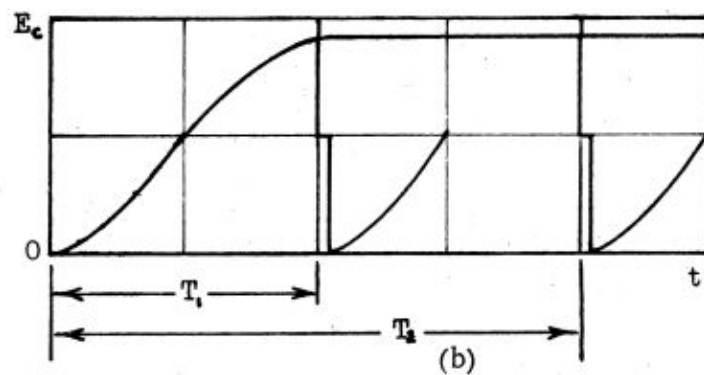
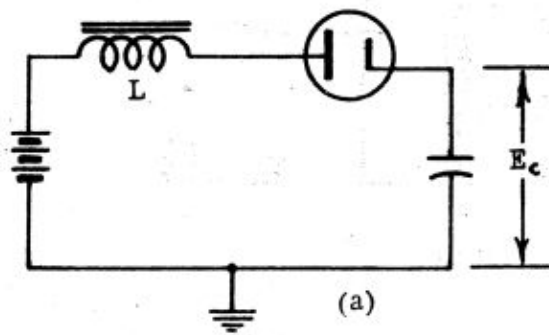


Fig. 12-13

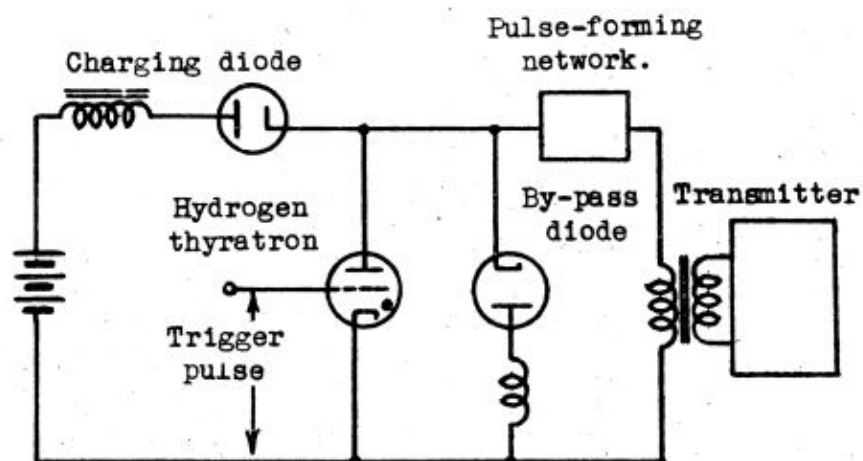


Fig. 12-14

RADIO NAVIGATIONAL AIDS.

an oscillation in the circuit comprising L and the capacitance C presented by the pulse-forming network. The voltage across C consists of a d-c component equal to the supply voltage, and an a-c component at the resonant frequency, having a peak value equal to the d-c component. Initially the a-c component is at its negative peak, as in Fig. 12-12 (b), so that the total voltage across C is zero. Half a cycle later the a-c component is at its positive peak, and the voltage across C is double the supply voltage. The spark gap now fires, and the pulse-forming network is rapidly discharged through the primary of the load transformer.

The charging period is controlled by the values of L and C ; it is half a cycle of the free oscillation, and so is given by $T = \pi\sqrt{LC}$ sec. This determines the repetition frequency, since the discharge is rapid in comparison.

If the gap did not fire, the oscillation would continue on as shown dotted. The circuit resistance is by no means negligible, and the oscillation is damped out in a few cycles.

A diode may be added in series with L as shown in Fig. 12-13 (a). The current is then prevented from reversing, and the system remains charged. It may be discharged at any later time desired, as indicated in Fig. 12-13 (b). The pulse repetition frequency may be lowered in this way, but not raised.

The use of a hydrogen thyatron for discharge of a line may introduce a special difficulty. Such a device conducts in one direction only, whereas all the other discharge switches conduct in either direction. A mismatch of the load to the line may cause the voltage across the line on the arrival of the first reflected wave to be negative. The thyatron will then stop conducting, and the negative voltage remains. This has the same effect as increasing the supply voltage for the next charging period, and a greater negative voltage will be left after discharge. The voltage will keep on increasing with every cycle.

Thus in Fig. 12-14 the pulse network is charged up with the left side positive, but this may become negative after discharge. The voltage in the charging circuit is then increased. The negative voltage may be prevented by means of a by-pass diode, which would short-circuit such a negative voltage but would have no effect on normal operation.

12-13 A-c Diode Charging.

A pulse-forming network may be charged from an a-c supply through a diode, as in Fig. 12-15 (a). The voltage builds up to the maximum, following the applied voltage, and remains at this value until the discharge impulse, as shown in Fig. 12-15 (b). The discharge must take place while the supply voltage opposes conduction, otherwise the supply would be short-circuited through the diode and discharge device. The pulse repetition frequency is either equal to the supply frequency, or a submultiple, such as one-half or one-third. The triggering pulses must be derived from the supply in some way, to maintain the necessary synchronism.

12-14 A-c Resonance Charging.

The circuit of Fig. 12-16 is convenient when the voltage required is high,

because the supply voltage needed is less than with other circuits. The pulse forming network may be represented by its input capacitance C_{in} , and an inductance L is connected in series with it to resonate at the frequency of the applied a-c voltage. Under these conditions the amount of energy supplied to the tuned circuit is at first greater than that dissipated, and the excess goes to build up the oscillation. Fig. 12-17 shows the condenser voltage for the case when the circuit has a magnification factor of 10. The first positive peak is about 2.7 times the peak value of the applied voltage. For a very large magnification factor, a multiplication of π is obtained instead of 2.7.

Fig. 12-18 shows the wave-form obtained if the pulse network is discharged at the first positive peak. The oscillation immediately starts to build up again and all cycles are the same.

12-15 Rotary Gap Modulator.

Fig. 12-19 shows the circuit of a typical rotary-gap modulator. The d-c voltage is obtained from a voltage-doubler rectifier, using two diodes. The only filtering provided is that through the two condensers charged by the respective diodes. The charging of the pulse-forming network is by d-c resonance with a diode. The spark wheel has four movable contacts set at right angles, and there are two fixed contacts at 45° . When both are used there are eight sparks per revolution, but if one is switched out of circuit there are only four. The output is taken through a 50-ohm cable to a pulse transformer. The transmitting valve, called a magnetron, has its anode earthed and negative H. T. pulses are applied to the cathode from the pulse transformer secondary. Two wires are wound together with the same number of turns, so that the two windings are closely coupled and have equal inductance. By-pass condensers are connected between the windings at both ends. The two windings are connected in the heater circuit, one in each lead. In this way the applied H. T. has no effect on the heater voltage. The condensers help here by by-passing the pulse frequency but not the heater frequency. A spark gap is connected across one secondary winding as a protection in case the valve fails to oscillate. An overload relay, normally closed, has its contacts connected in the a-c supply lead, and the relay coil carries the charging current of the pulse forming network. In the event of overload the relay opens and cuts off the power.

Portion of the output is taken by means of a voltage divider following the pulse-forming network, to serve as a triggering pulse for the timing circuits, which must be of flip-flop type.

12-16 Oscillations after Pulse.

The secondary circuit of the pulse transformer contains leakage inductance and also capacitance due to winding, wiring, and the valve. It is damped mainly by the transmitting valve. The sudden changes at the beginnings and end of each pulse tend to produce r-f oscillations in this circuit, as in Fig. 12-20 (a).

The oscillation due to the leading edge of the pulse is heavily damped by

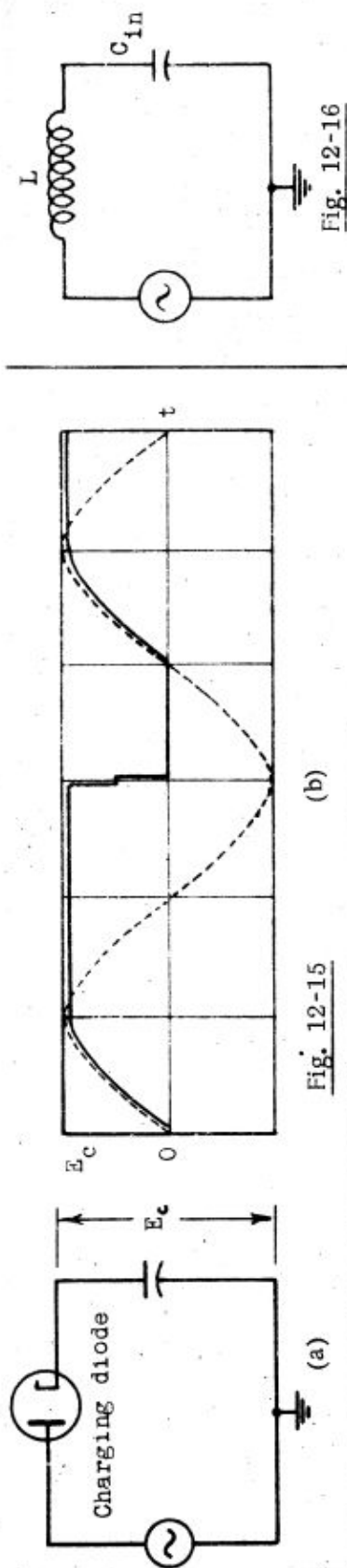


Fig. 12-16

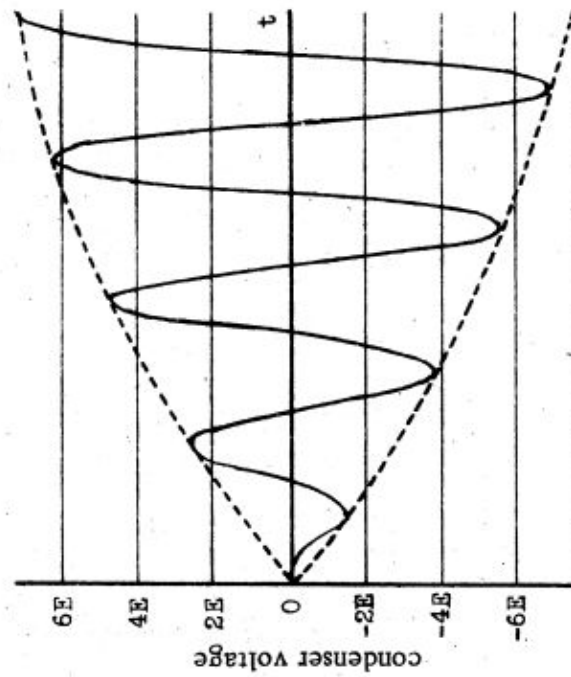
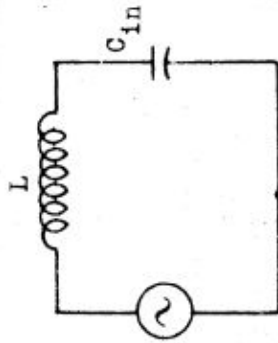


Fig. 12-17

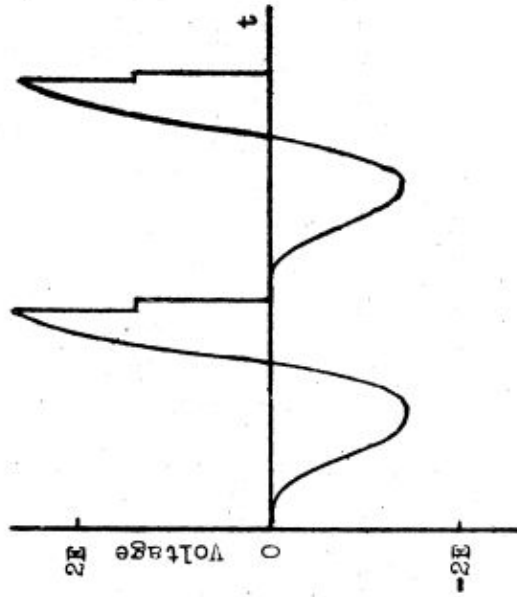


Fig. 12-18

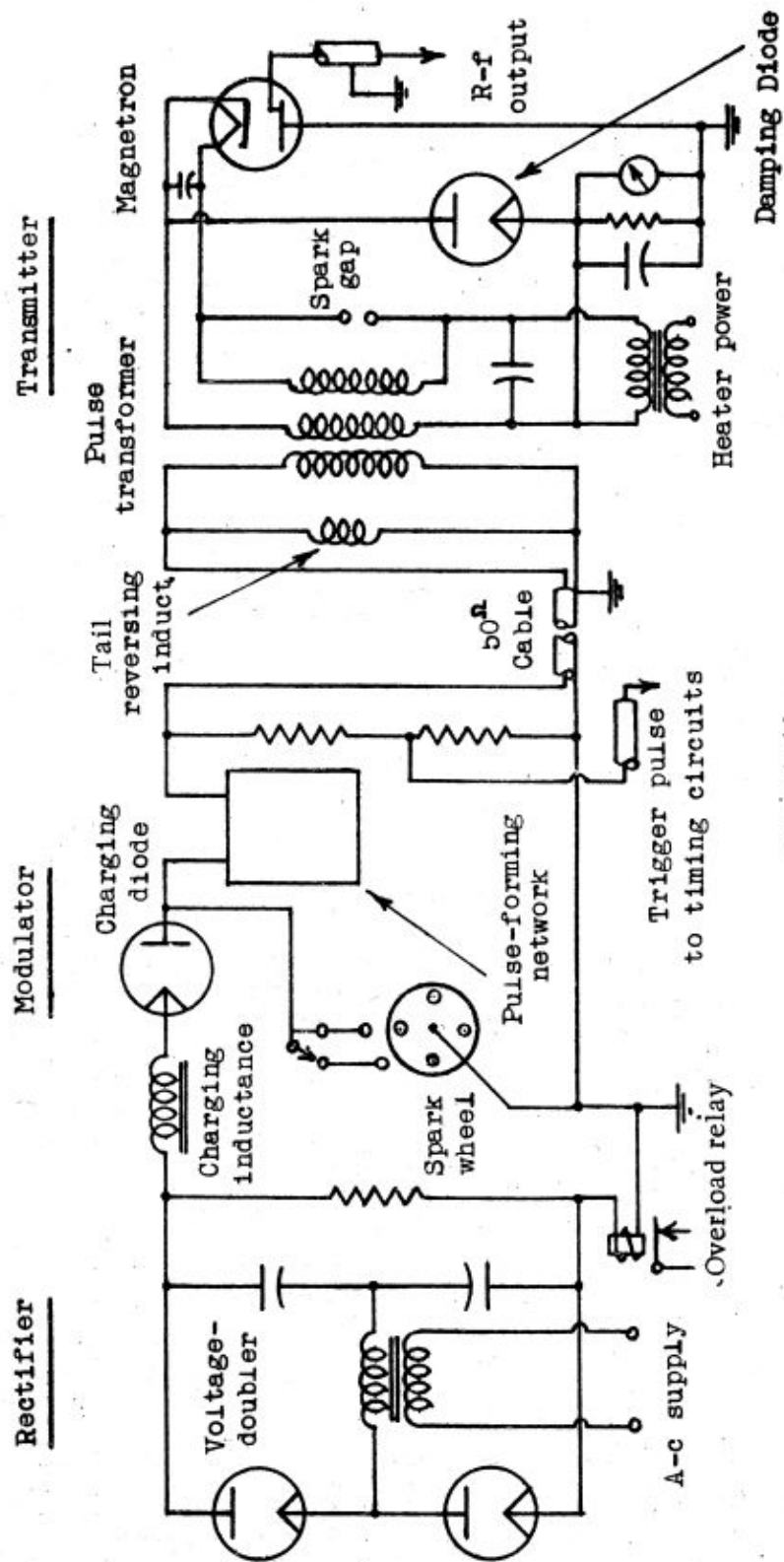


Fig. 12-19

RADIO NAVIGATIONAL AIDS

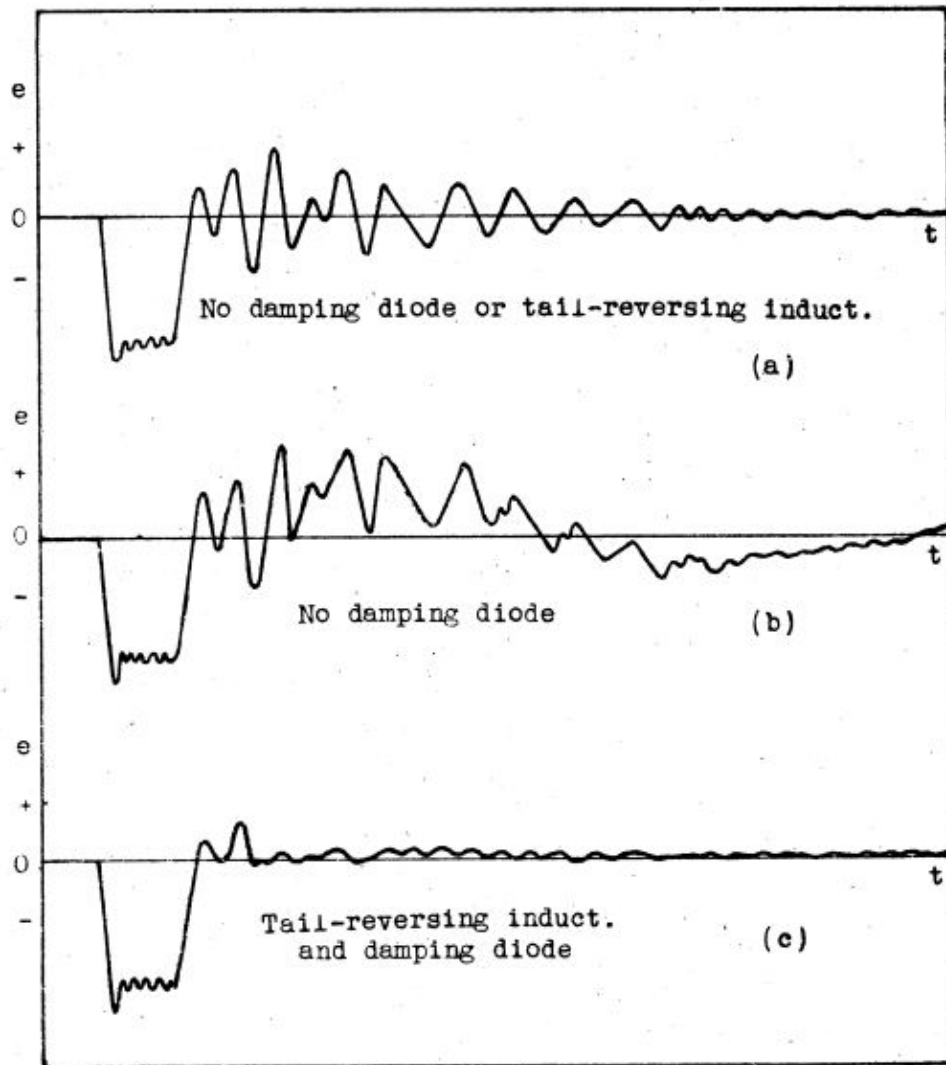


Fig. 12-20

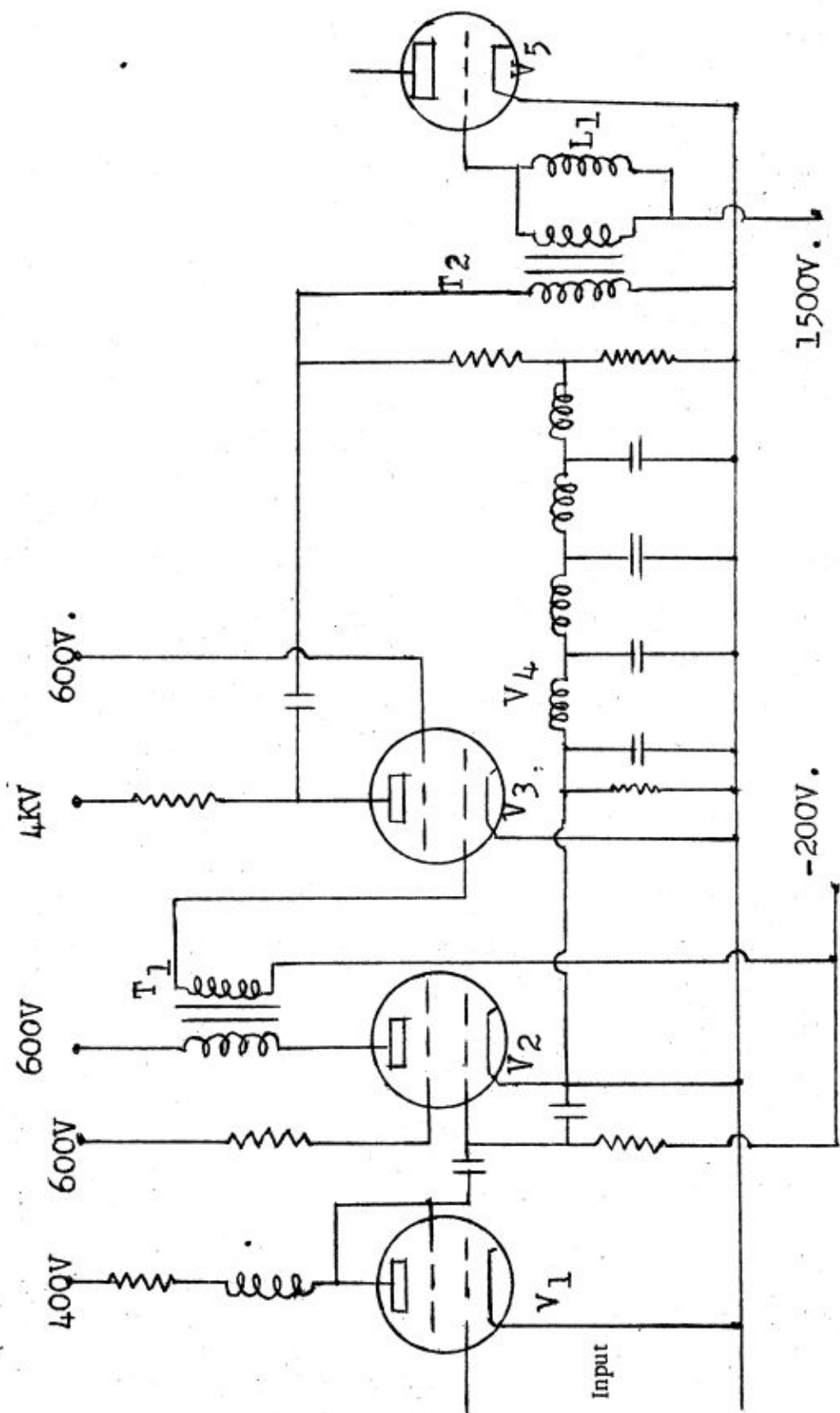


Fig. 12-21

RADIO NAVIGATIONAL AIDS.

the conducting valve, and does not build up any great amplitude. But after the end of the pulse the H. T. voltage is removed, and so the damping is much smaller. The oscillation can build up to a large value and persist for a time much longer than the pulse length. On the negative peaks the valve may oscillate, and give the false appearance of a reflection.

This trouble may be cured by means of a "tail-reversing" inductance connected across the pulse transformer primary, in combination with a diode across the secondary, so connected that it conducts when the H. T. voltage of the transmitting valve is reversed. The diode is therefore idle during pulses. The current in the tail-reversing inductance builds up during pulses, and causes a low-frequency oscillation to be produced immediately after the pulse, as the current continues to flow through the spark gap into the pulse network. The effect of this oscillation is to make the secondary voltage more positive during the period immediately following the pulse, but more negative later, changing Fig. 12-20 (a) to Fig. 12-20 (b). When the secondary voltage is positive the diode conducts and damps out the oscillation, as in Fig. 12-20 (c).

12-17 Pulse Transformers.

Instead of using a very high supply voltage, a moderate voltage may be used, and the final voltage obtained by passing the pulse through a special pulse transformer.

In order to pass a pulse without appreciable distortion, the transformer must be different in design from an ordinary power-frequency transformer. The design is based on measurements of the magnetic behaviour of the iron for short pulses. The essential point is to keep the leakage inductance and the self-capacitance as small as possible, but the primary inductance need not be very high. Close coupling between primary and secondary is important. The iron laminations are very thin, about 3 mils (.003 inch) being common. Relatively few turns are used, and the voltage per turn is high. The wire is heavily insulated, and insulating barriers are placed between layers. The whole transformer is oil-immersed, and it is common practice to seal the case, and to provide a flexible diaphragm to allow for oil expansion due to temperature changes. Very good insulation must be provided between the separate windings, and between windings and core.

12-18 Delay-line-controlled "hard-valve" modulator.

An entirely different type of modulator to the one previously discussed is shown in Fig. 12-21. This consists of a driver/power-amplifier combination designed to produce the required high voltage pulse. The final power amplifier valve, which is called upon to produce a pulse of the order of 4-5kv in amplitude, will require to be a large valve, and frequently several valves are connected in parallel to obtain the required results. The driver stages will be required to deliver some considerable grid power to the final amplifier stage.

In Fig. 12-21, V_1 , a 6L6, and V_2 , an 829B, form the first two stages of the driver, V_3 and V_4 , two 829B's are connected in parallel to form the final

driver stage. The final amplifier stage V_5 , although shown in the diagram as one valve for the sake of simplicity, would probably consist of 2, 3, or even 4 valves type 6C21 or similar all connected in parallel.

In order to study the operation of the circuit, consider the normal condition, when V_1 is conducting, and V_2 , V_3 , and V_4 are cut off. To pulse the transmitter, a negative pulse, Fig. 12-22 (a), is applied to the grid of V_1 from a timing circuit, cutting off the valve. Typical values of pulse duration are shown in Fig. 12-22, although these would of course be varied according to the particular application. The resulting positive pulse at V_1 anode, Fig. 12-22 (b), is coupled to the grids of V_2 , producing a negative pulse at V_2 anode. This negative pulse is coupled via a pulse transformer T_1 to the grids of V_3 and V_4 , being stepped up in the transformer, and applied positive going at the grids of the succeeding stage, Fig. 12-22 (c). A negative pulse at the anodes of V_3 and V_4 , Fig. 12-22 (d), is coupled via a second pulse transformer T_2 to the grid of V_5 , the final power amplifier stage, Fig. 12-22 (e).

The output pulse, in this case 0.8 microseconds, is shorter than the original input pulse from the timing circuit, because the voltage at the primary of T_2 is fed back through a potential divider and a 0.8 microsecond delay line to the grid of V_2 . The negative going leading edges of the delayed pulse cuts off V_2 before the end of the input pulse occurs. Because the leading edge of the input pulse from the timing circuit controls both edges of the final pulse, it must be steep, and furthermore, the pulse must hold V_1 at cut off for at least 0.8 microseconds in order to preserve the flat top on the final pulse. V_1 must return to conduction shortly after 0.8 microseconds as the feedback pulse is coupled to V_2 grid through a short time-constant circuit, and can only hold V_2 cut off for a very short time.

The positive going wavefront applied to the L. H. end of the delay line at the commencement of the pulse will not effect the operating of the circuit. The wavefront will travel to the R. H. end of the line, where the line is terminated in its correct impedance, and negligible reflection will take place.

At the completion of the pulse formation, a positive going wavefront Fig. 12-22 (d), is applied to the R. H. end of the line. The wavefront reaches V_2 grid 0.8 microseconds after the end of the pulse (the 0.8 microsecond pulse which has just been completed) and would tend to cause a second conduction in V_2 . Because the resistance of V_1 when conducting is much less than the characteristic impedance of the delay line, the increase of V_2 grid voltage is too small to cause conduction.

A tail-sharpening inductance L_1 is connected across T_2 secondary, as without it, the trailing edge of the input of V_5 would not be very steep. The application of the inductance increases the rate of discharge of shunt capacitances, because of current built up in the inductance during the pulse.

12-19 The Magnetic Modulator.

It will be recalled that the function of the modulator is to supply recurrent pulses of power to the transmitter. These pulses are required to occur at a fixed

RADIO NAVIGATIONAL AIDS

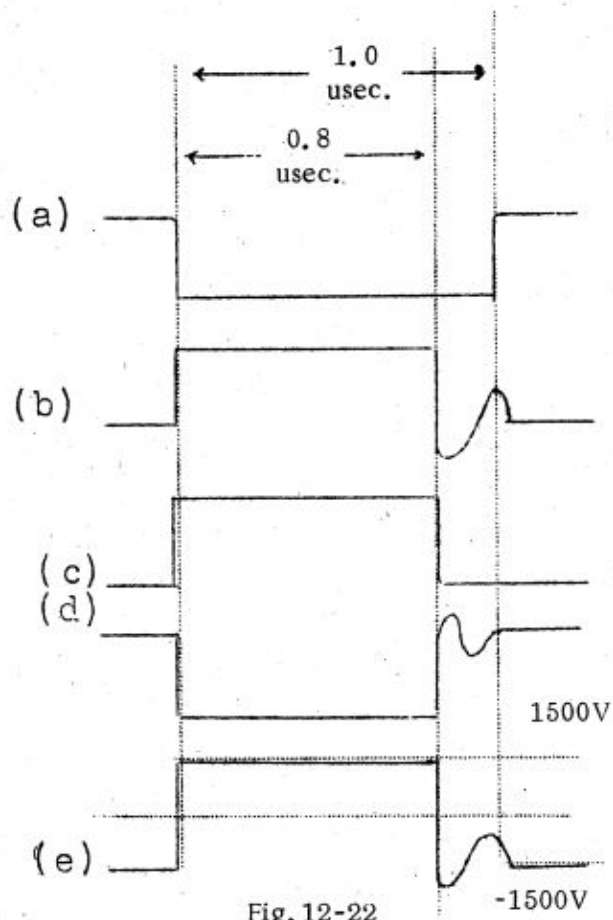


Fig. 12-22

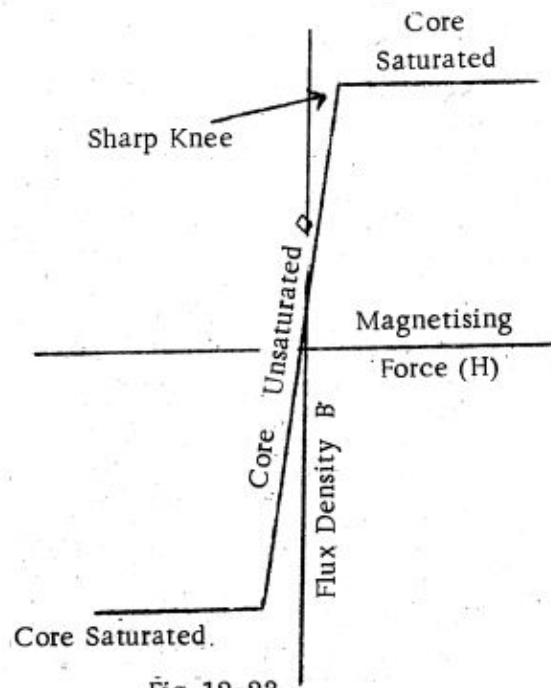


Fig. 12-23

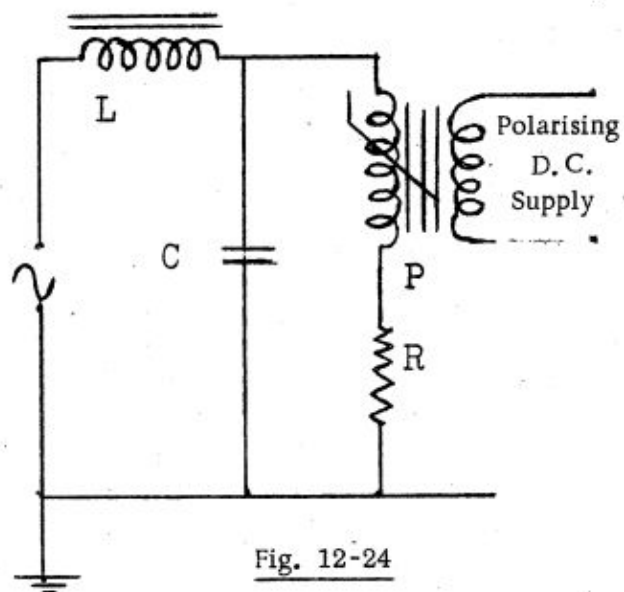


Fig. 12-24

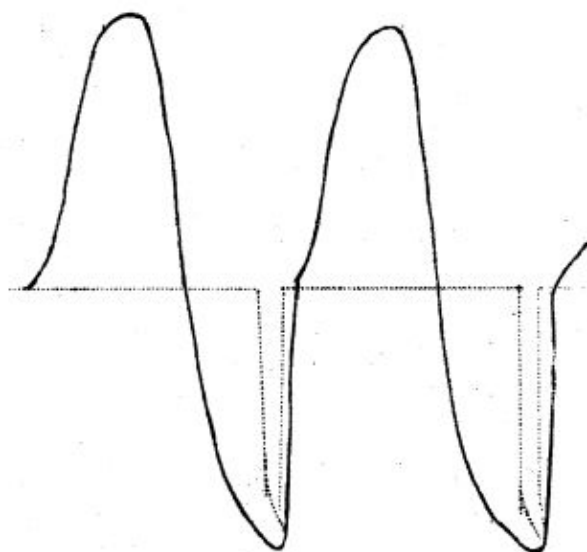


Fig. 12-25

RADIO NAVIGATIONAL AIDS.

repetition rate, to be of fixed duration, and to be of sufficient amplitude to operate the transmitter. The pulses are rectangular in form, and are generally produced by the discharge of a previously charged pulse forming network through a suitable switch into a matched load. The pulse forming network consists of a number of sections of shunt capacitance and series inductance arranged to simulate the distributed capacity and inductance of an open-circuited transmission line.

12-20 The Pulsactor.

In the magnetic modulator, magnetic switches are used both for discharging the pulse forming network and also in the charging circuit. These switches are known as pulsactors, and are a special case of the saturable reactor. An inductance may be designed so that its core saturates at comparatively low coil currents. Such inductances use cores made of special materials such as "Mumetal" which possess a rectangular type of hysteresis loop. The special characteristics of such inductances are a very high permeability when the core material is in the unsaturated state, and a very low permeability when the core material is saturated. In other words, when the core is not saturated, the inductance and impedance are high, but when the core is saturated, the inductance and the impedance are low. Transition from one condition to the other takes place very rapidly, as indicated by the sharp "knee" of the B-H characteristics of a typical core material, as shown in Fig. 12-23. As the inductance of a reactor is proportional to the permeability of its core material, it will be seen that, when saturation of the core takes place, there is a sudden drop in the inductance of the pulsactor, and a corresponding drop in the pulsactor impedance. If now we examine the essential features of a conventional switch, we see that they are very similar to the pulsactor in that, whilst the switch is "open", it has an extremely high impedance, whereas when it is "closed", its impedance is very low. The pulsactor, having similar properties, may be said to have a switching action.

Since the device is neither electronic nor mechanical, it may be expected to have a long satisfactory life. In the past however, a number of disadvantages have militated against a very widespread use of the pulsactor, particularly in high range discrimination equipments such as marine navigational apparatus. These disadvantages are (a) high shunt capacitance, (b) relatively low efficiency of the associated charging circuit (c) pulse shapes obtained were poorer than those obtained with other switching devices (d) core materials available in the past have been unsuitable for the production of very short pulses. Now with the production of new materials and improved techniques, the magnetic modulator has become a much more practical proposition, and is being used in navigational equipment.

12-21 Pulsactor Circuits.

To make use of the switching properties of pulsactors, they must be incorporated in specially designed circuits, in which saturation of the pulsactor core

is arranged to take place as maximum voltage on a storage condenser or other storage device is realised. Turning to an examination of Fig. 12-24, which shows one basic circuit, L and C are chosen so as to be resonant at the supply frequency, and the inductance of the unsaturated reactor P is designed to be so high that it has very little effect on the charging of C . P is polarised by means of an external D.C. source, that is, moved up the sloping portion of the hysteresis curve in Fig. 12-23, so that it will saturate only in one direction, and the pulses obtained will be unidirectional. The waveform of condenser voltage V_c is shown in Fig. 12-25. As the circuit is resonant, V_c is 90° out of phase with the supply voltage, and at the end of one complete cycle, the supply voltage is zero, while V_c has reached a maximum. The design of P is such that its core saturates at this point, and C discharges through the low inductance of the saturated pulsactor, the charge being transferred into the load R in the form of a pulse, the character of which is determined by the values of C , R and L_s , the inductance of the saturated pulsactor. V_c having now been reduced to zero, conditions in the circuit at the end of the pulse are exactly as they were at the beginning of the cycle of operations, so that pulse formation will take place recurrently at the supply frequency.

In this circuit the properties of the pulsactor are used in two ways. Firstly, the device is used in the unsaturated condition to hold off the condenser charging voltage from the load resistance, until the charging voltage reaches a maximum. Then, secondly, the device is used in the saturated condition to discharge the condenser, allowing the load to receive energy stored in the condenser. It will be appreciated that there must be one part of a modulator circuit which is concerned with the charging of the pulse forming network, and one part which is concerned with the discharge. If we have a pulsactor switch to discharge the pulse forming network, then in order that the pulse may remain substantially rectangular and its shape be unchanged by the presence of the pulsactor, the inductance of the saturated pulsactor must be very small, and comparable with one of the inductances forming the pulse forming network. The discharge circuit will then be satisfactory. A further requirement is that of the charging circuit, that is, that the charging of the pulse forming network must be substantially unaffected by the presence of the pulsactor, whose inductance when unsaturated must therefore be very high.

12-22 The Pulsactor Cascade Discharge Circuit.

Frequently a single pulsactor switch cannot meet these two requirements of the charge and discharge circuits, so a cascade connection of two or more pulsactors is used. To study the operation of such a circuit, Fig. 12-26 shows the basic arrangement, and the design requirements are:-

- (1) $C_1 = C_2 = C_3 = C_n$.
- (2) Unsaturated impedance of P_1 is high enough to allow C_1 to be charged.

RADIO NAVIGATIONAL AIDS

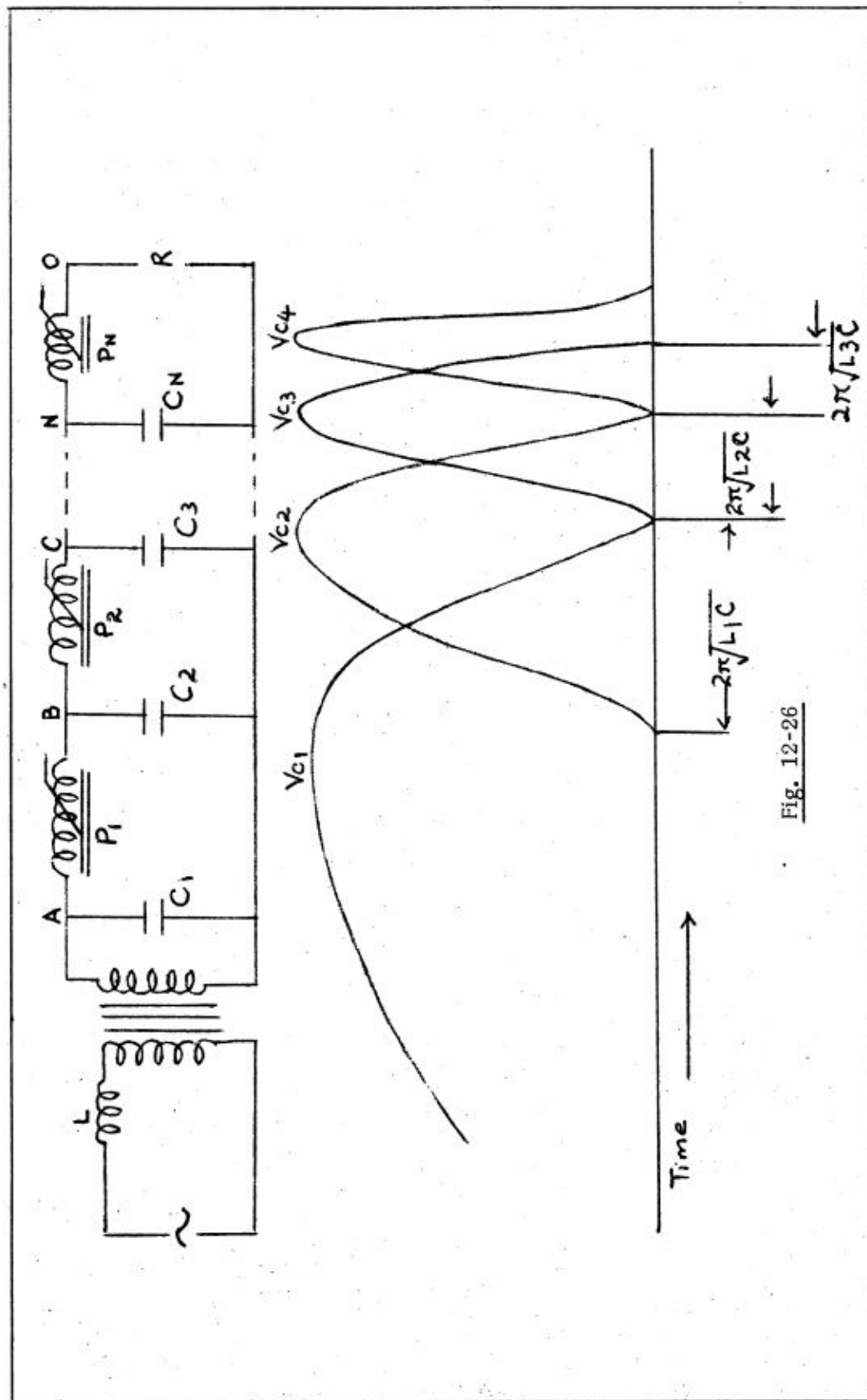


Fig. 12-26

- (3) Unsaturated impedance of P1 much greater than the unsaturated impedance of P2 which must be much greater than the saturated impedance of P1.
- (4) Unsaturated impedance of P2 much greater than the unsaturated impedance of P3, which must be greater than the saturated impedance of P2. Etc. Etc.
- (5) Unsaturated impedance of Pn much greater than the load impedance R, which must be much greater than the saturated impedance of Ph.
- (6) Linear choke L, step up transformer T, and condenser C are resonant at the supply frequency.

Fig. 12-26 (a) shows the basic circuit arrangement of such a system, and Fig. 12-26 (b) shows the voltage waveforms on the various condensers in the circuit. The time scale in Fig. 12-26 (b) is very much expanded compared with that of Fig. 12-25, and the waveform Vc1 in Fig. 12-26 (b) is practically the same as waveform Vc in Fig. 12-25, but to a different time scale. While C1 is being charged, point B is held almost at earth potential, since the impedance of P2, P3, etc. is much less than the impedance of P1, and therefore the major portion of Vc1 appears across P1. All the reactors are still unsaturated. The design of P1 is such that it becomes saturated when Vc1 is at a maximum, causing a very large reduction in the impedance of P1. This in turn results in an oscillatory discharge of C1 into C2 through the saturated impedance of P1. The period of this discharge is

$2\pi\sqrt{L_1C}$, where L_1 represents the saturated impedance of P1, and C represents the effective capacitance of C1 and C2 in series. Vc2 rises as shown in Fig. 12-26 (b), and after half a period of the oscillation reaches a maximum, which is equal to the peak voltage which appeared across C1 before the discharge took place. The current in P1 has now fallen to zero, and P1 becomes once again a high impedance, preventing the rapid discharge of C2 back into C1. In a similar manner, when Vc2 reaches a maximum P2 saturates, and the charge is transferred from C2 to C3 in an even shorter time, since the period is now $2\pi\sqrt{L_2C}$. L_2 represents the saturated impedance of P2, and C being as before, the effective capacitance of C2 and C3 in series. L_2 , following the design requirements previously stated, is much less than L_1 .

This cascading of pulsactors and condensers can be continued as far as desired, until the losses in the pulsactors, which increase with the increasing rates of change of flux, become so large that the design requirements cannot be met. Each condenser charges at a more rapid rate than the previous one, and the last condenser is switched on to the load R as in the case of single stage circuit as shown in Fig. 12-24. In the practical magnetic modulator, the final condenser of the cascade system is replaced by a pulse forming network, which charges in the same way as would a condenser of the same total capacitance, but gives a rectangular pulse on rapid discharge.

RADIO NAVIGATIONAL AIDS

QUESTION PAPER FOR CHAPTER 12.

1. Explain the effects of mis-match between a pulse-forming line and the resistance through which it is charged.
2. Explain the action of the grid in a thyatron.
3. What are the disadvantages of a simple resistance - capacitance pulse forming network ?
4. Describe the AC resonance method of charging.
5. What means can be used to suppress oscillations after a pulse ?
6. How is a delay line used to control the duration of the output pulse of a hard valve modulator.
7. Explain the use of a pulsactor as a discharge device.

- - - 0 - - -



A Service of
AMALGAMATED WIRELESS (AUSTRALASIA) LTD.

